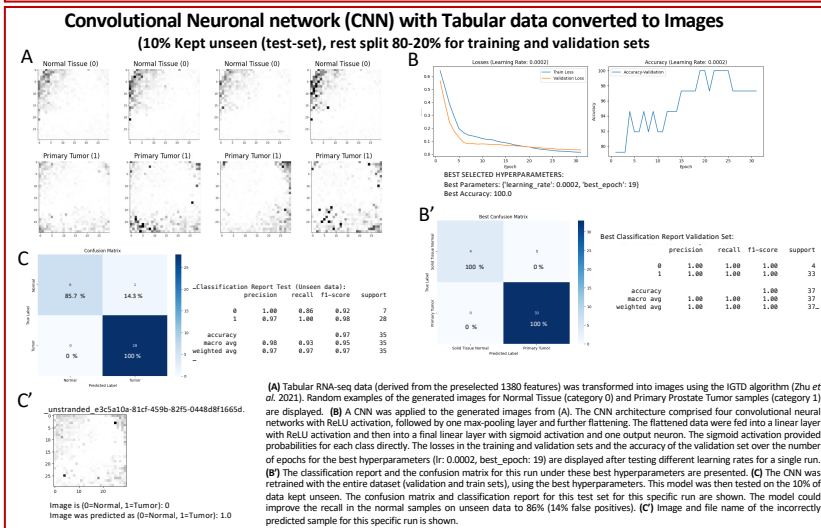
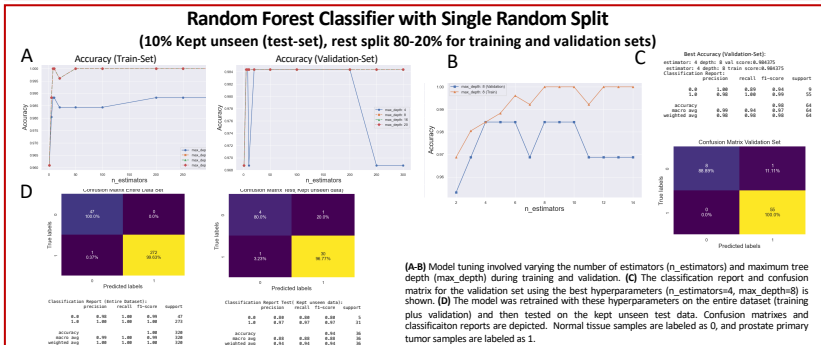
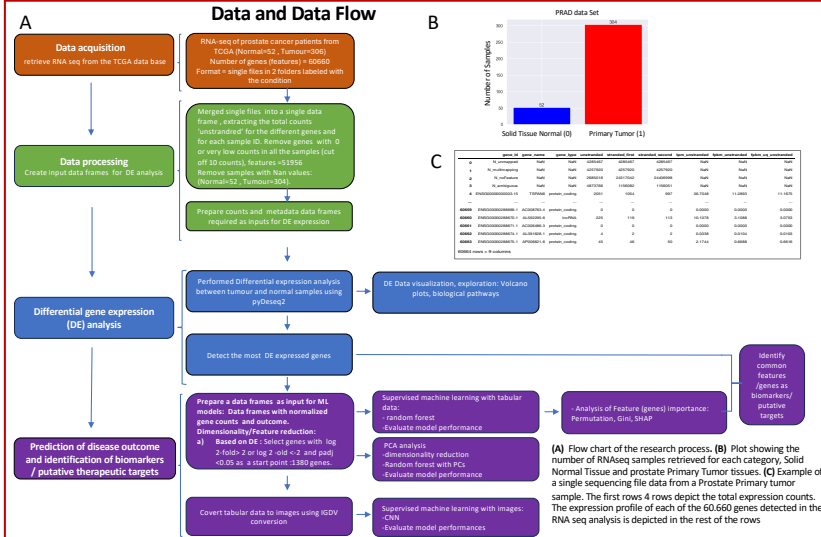


Prostate Cancer Prediction and Biomarker Identification Using Machine Learning and Deep Learning Algorithms on Transcriptome Data from The Cancer Genome Atlas (TCGA) Database

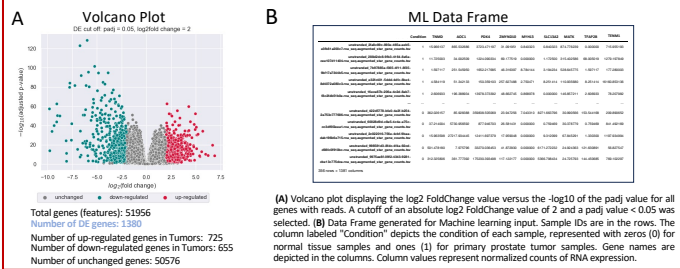
ABSTRACT

The search for novel RNA biomarkers and innovative methods to identify cancerous tissues can significantly advance the development of RNA-based diagnostic and therapeutic strategies, leading to more effective and personalized approaches for cancer treatment and management. In this project, we investigated the feasibility of predicting or diagnosing prostate cancer, which ranks among the most prevalent cancers in the male population, by applying machine learning (ML) and convolutional neural network (CNN) algorithms to gene expression data of normal and primary tumor prostate gland samples. Genes/features used as input for ML were reduced by preselecting the most differentially expressed (DE) genes between cancer and normal samples. Machine learning algorithms (logistic regression, random forest, random forest on the most important principal components (PCs)) were applied to predict cancer outcomes using RNA expression data on the selected genes. A CNN was also tested on the same tabular data converted to images. Moreover, through an examination of the disturbed gene expression patterns in prostate cancer samples and the genes important for predicting cancer versus normal tissue outcomes by machine learning, we also set up to discover putative novel RNA biomarkers for prostate cancer.

RESULTS

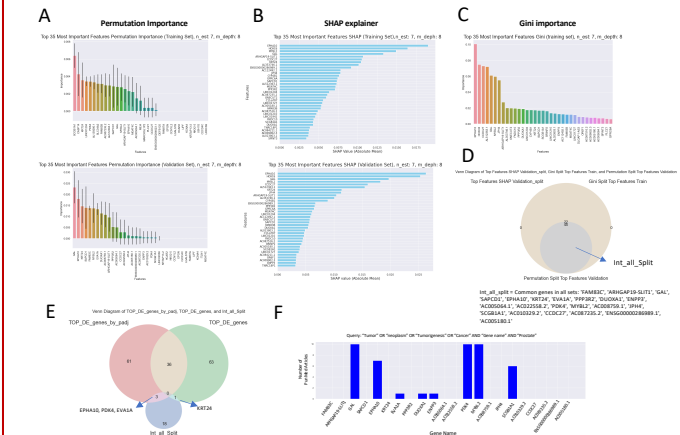


Differential Expression (DE) Analysis and Machine learning Data preparation



Performance of the Different Models Tested						
Data / Features	Model	Training and evaluation	Accuracy Validation set	Recall Validation set	Accuracy Test set (unseen data)	Recall Test set (unseen data)
- Tabular data - Features reduction from DE analysis log2 fold > 2 and log2 < -2 and padj < 0.05 - 1380 genes	Logistic Regression	35 iterations	0.931	0.82(0), 0.95 (1)		
		50 iterations (keep 30% unseen)	0.953	0.78(0), 0.98(1)	0.94	0.6 (0), 1(1)
		Grid-Search-single-random-split	0.986	0.91 (0), 1 (1)		
	Random Forest	Grid-Search-single-random-split (keep 30% unseen)	0.984	0.89 (0), 1 (1)	0.94	0.8 (0), 0.97 (1)
		Crossvalidation	0.955 (mean CV score) 0.971 (median CV score)	0.8 (0), 1(1) (median)		
		Crossvalidation (keep 30% unseen)	0.956 (mean CV score) 0.965 (median CV score)	0.8 (0), 1(1) (median)	0.94	0.8 (0), 0.97 (1)
- Tabular data - Features reduction by DE follow PCA analysis - Most important PCs used for ML - PCA1,2	Random Forest	Grid-Search-single-random-split	0.972	0.91(0), 0.98(1)		
		Grid-Search-single-random-split (keep 30% unseen)	0.984	0.89(0), 1(1)	0.94	0.8 (0), 0.97 (1)
		Crossvalidation	0.958 (mean CV score) 0.971 (median CV score)	0.82(0), 0.97(1) (median)		
	CNN	Grid-Search (keep 30% unseen)	0.967 +/- 0.018	0.945 +/- 0.071 (0), 0.971 +/- 0.071 (1) (mean 6 runs)		
		Grid-Search-4CINs-ReLU activation	0.874 +/- 0.054 (mean 6 runs)	0.147 +/- 0.3 (0), 0.997 +/- 0.008 (1) (mean 6 runs)		
		Grid-Search-4CINs-ReLU activation (keep 10% unseen)	0.978 +/- 0.021 (mean 6 runs)	0.908 +/- 0.089 (0), 0.978 +/- 0.021 (1) (mean 5 runs)	0.952 +/- 0.016 (mean 5 runs)	0.886 +/- 0.072 (0), 0.974 +/- 0.021 (1) (mean 5 runs)

Feature Importance Analysis to identify Biomarkers or Therapeutic Targets



CONCLUSIONS AND OUTLOOK

- Machine learning applied to RNAseq data has successfully predicted prostate cancer outcomes.
- Random forest outperformed logistic regression, enhancing recall for under-represented normal tissue.
- PCA feature reduction was effective; 2 PCs matched RF performance with 1,380 features.
- Transforming tabular data into images for a CNN improved model performance, particularly recall for the underrepresented category; visualization provided insights not easily discernible from 1,380 tabular features.
- Main issues: unbalanced, limited data and no accessible independent dataset for final validation. While models showed high accuracy, they struggled with underrepresented normal samples but excelled in classifying tumor samples.
- Stratified splitting improved Random Forest performance on underrepresented samples. Further enhancement of CNN could be achieved with stratification and cross-validation.
- Generating synthetic RNA-seq data and utilizing independent datasets is recommended.
- Optimizing Random Forest by adjusting hyperparameters (min_samples_split, min_samples_leaf, max_features) is advised to boost stability, reduce overfitting, and enhance performance.