

Enhancing Risk Management in Financial Institutions through Machine Learning and Natural Language Processing: A Critical Evaluation of AI Techniques

MSc Research Project
MSCFTD1

Anjaly Elias
Student ID: X23302925

School of Computing
National College of Ireland

Supervisor: Victor Del Rosal

National College of Ireland
MSc Project Submission Sheet
School of Computing



Student Name: Anjaly Elias.....

Student ID:X23302925.....

Programme:...MSCFTD1..... **Year:** ...2024-2025.....
 :

Module: ...Practicum.....

Supervisor: ... Victor Del Rosal.....

Submission

Due Date: ...15/09/2025.....

Project Title: ENHANCING RISK MANAGEMENT IN FINANCIAL INSTITUTIONS THROUGH MACHINE LEARNING AND NATURAL LANGUAGE PROCESSING:A CRITICAL EVALUATION OF AI TECHNIQUES.

Word

Count: ...9539..... **Page Count:**.....24.....

I hereby certify that the information contained in this (my submission) is information pertaining to research I conducted for this project. All information other than my own contribution will be fully referenced and listed in the relevant bibliography section at the rear of the project.

ALL internet material must be referenced in the bibliography section. Students are required to use the Referencing Standard specified in the report template. To use other author's written or electronic work is illegal (plagiarism) and may result in disciplinary action.

Signature:Anjaly.....

Date:15/08/2025.....

PLEASE READ THE FOLLOWING INSTRUCTIONS AND CHECKLIST

Attach a completed copy of this sheet to each project (including multiple copies)	☐
Attach a Moodle submission receipt of the online project submission , to each project (including multiple copies).	☐
You must ensure that you retain a HARD COPY of the project , both for your own reference and in case a project is lost or mislaid. It is not sufficient to keep a copy on computer.	☐

Assignments that are submitted to the Programme Coordinator Office must be placed into the assignment box located outside the office.

Office Use Only	
Signature:	
Date:	
Penalty Applied (if applicable):	

Enhancing Risk Management in Financial Institutions through Machine Learning and Natural Language Processing: A Critical Evaluation of AI Techniques

Anjaly Elias
X23302925

Abstract

The financial services industry is undergoing a significant transformation driven by the integration of artificial intelligence. Effective risk management, a cornerstone of institutional stability, stands to benefit immensely from these technological advancements. This thesis explores the application of machine learning (ML) and natural language processing (NLP) to enhance credit risk assessment, a critical component of financial risk management. The whole Lending Club loan data set will be used in establishing a reliable baseline for modelling, following a robust data pre-processing, balancing, and feature engineering process. This research will evaluate the performance of three distinct AI methods- an ensemble method from the historical world (Random Forest), a probabilistic classifier (Gaussian Naive Bayes), and a deep learning model (Long Short Term Memory network). All of those models make use of a balanced dataset that consists of 100,000 loan applications. The approaches are rigorously compared by performance measurement of accuracy, precision, recall, F1 score, and AUC. The results show that LSTM models achieved the highest predictive accuracy of 73.55%, improving the tuned random forest performance of 73.14% and outpacing significantly naive Bayes models with 70.23%. This shows that deep learning architectures have the potential to capture complex, non linear patterns in financial data that can be missed by traditional models. The last phase of research involves design creation and implementation of a prototype web application demonstrating a real avenue through which advanced analytic might be deployed in a financial institution's workflow. Findings point towards the promise of AI to upgrade risk predictions models and painstaking data preparation and evaluation processes.

1 Introduction

An ever-compounding element of the dynamic ambience of global finance is the pertaining effective risk assessment, which, besides being a regulatory prerequisite, is the substratum for survival as successful existence of any institution. Behaviour by a financial institution, interacting with a web of risks, would include market risk, operational risk, and credit risk,

amongst others. Not limited to, but very much foremost in the comprehensive realm of concerns, is credit risk, posed herein as the risk of loss due to a borrower's inability to repay the loan (Dewasiri et al., 2024). Conventional techniques taken for credit risk appraisal were predominantly statistical models, such as logistic regression, utilizing a constricted array of financial parameters, including credit score, income levels, etc. Although these techniques have nearly been the gold standards for decades, they have a habit of inadequately assessing the borrower's financial profile with changing circumstances and at a much lower speed.

The digital age has triggered an avalanche of data with a concurrent revolution of computational capabilities, favoring the adoption of artificial intelligence (AI) and machine learning (ML) for application in finance. Such modern technologies present an opportunity to overcome the limitation of traditional models by sifting through huge, heterogeneous datasets to identify subtle patterns and non-linear relationships pertaining to risk assessment (Pattnaik et al., 2024). The AI and ML technologies are becoming indispensable between newer horizons with the growth of FinTech in algorithmic trading, fraud detection, and regulatory compliance (Kanaparthi, 2024; Addy et al., 2024). As these systems can learn from data and improve with time, this represents a pathway of disruption for financial institutions to proactively manage and mitigate risk.

On the larger AI canvass, NLP is regarded as one among the most liberalizing technologies that have begun to sweep through the financial domain (Gao et al., 2021). Unstructured data accounts for a large chunk of financial information in text format over loan applications, news articles, analyst reports, and internal communications. NLP techniques make it possible for computers to read, understand, and extract useful insights from unstructured human language data (Chen et al., 2024). In this manner, NLP could deepen the qualitative input to quantitative risk models by analyzing loan purposes stated in an applicant's own words or sentiment analysis of news coverage around an industry, thereby increasing predictive power. However, even with a dreamlike opportunity in sight, operators of sophisticated techniques face their set of challenges. Issues of model interpretability, data privacy, regulatory scrutiny, and the possibility of algorithmic bias require serious deliberation (Vuković et al., 2025; Kute et al., 2021). For that reason, it is not just a matter of mechanistically implementing the algorithm. Rather, critical reflection must take precedence in order to test for efficacy, identify potential drawbacks, and ensure a responsible approach to its application.

In this context, the thesis addresses the implementation and comparative performance of various artificial intelligence techniques applied to the problem of credit risk assessment. The main research question guiding this study is: **How can machine learning and natural language processing be applied, evaluated, and ultimately integrated toward enhancing credit risk prediction for financial institutions?**

To answer this question, the following research objectives have been established:

1. To collect, preprocess, and balance a large scale, real world financial dataset representing both approved and rejected loan applications to create an unbiased training environment.
2. To apply feature engineering techniques, including rudimentary NLP on textual data, and use dimensionality reduction to prepare the data for modeling.
3. To implement and train a diverse set of models, including a traditional ensemble classifier (Random Forest), a simple probabilistic model (Naive Bayes), and a sequential deep learning model (Long Short Term Memory).
4. To conduct a critical and comprehensive evaluation of the models' performance using a suite of standard metrics and visualization techniques to identify the most effective approach.
5. To design and develop a prototype system that demonstrates how the final, optimized model can be deployed for practical use in a financial setting.

This research contributes to the scientific literature by providing a detailed, end to end case study on applying and comparing different classes of AI models to the problem of credit risk. Unlike many studies that focus on a narrow set of algorithms, this work contrasts traditional ML with deep learning and discusses the full lifecycle from data ingestion to a deployable prototype. The findings offer valuable insights for financial practitioners seeking to leverage AI for enhanced risk management and for academics exploring the frontiers of financial technology.

The remainder of this report is structured as follows. Chapter 2 provides a critical review of the existing academic literature on AI, ML, and NLP in financial risk management. Chapter 3 details the research methodology, including the data source, preprocessing steps, and the selection of models and evaluation metrics. Chapter 4 presents the design specification for the proposed system, outlining its architecture and components. Chapter 5 describes the implementation of the data processing pipeline, the model training process, and the development of the prototype web application. Chapter 6 presents a thorough evaluation and discussion of the experimental results. Finally, Chapter 7 concludes the report, summarizing the key findings, discussing the limitations of the study, and proposing meaningful directions for future research.

2 Related Work

Essentially, this particular chapter attempts to provide a review of the whole literature in academia regarding the use of artificial intelligence, machine learning, and natural language processing concerning financial risk management. The review contextualizes the current piece of research against the available literature, pinpoints important trends and developments within the body of research, and trends gaps that this thesis seeks to address.

2.1 The Rise of Artificial Intelligence in Financial Risk Management

Bringing artificial intelligence within financial services is indeed one of the most popular subjects of increased research interest by both academic and industry people. Recently published systematic literature reviews have reflected the fast-paced adoption of AI and ML in dealing with lingering issues in finance (Abdulla & Al-Alawi, 2024; Pattnaik et al., 2024). For a long time, risk management relied on effective statistical models with linear assumptions and a limited number of variables. However, with the current financial ecosystem having both high data volume and complexity, time-proven statistical methods are at times inadequate. AI has transformed the whole financial management study by enabling a more dynamic and data-driven approach to decision making, argues Zakaria et al. (2023).

Several scholars have been interested in the broad influence of AI across various financial functions. In this regard, Addy et al. (2024) present a critical review of ML applications in algorithm-based trading as well as in risk management, and they conclude that machine-learning models consistently outperform their traditional counterparts regarding the identification of complex market patterns. On their part, Tian et al. (2024) went ahead to carry out systematic reviews of ML in internet financial risk management and discovered a strong leaning towards the adoption of ensemble and deep learning methods in managing the unique risks of online financial platforms. One of the key themes noticeable from this collection of works is the new trend from reactive to proactive risk management whereby AI models predict possible risks even before they occur and thus allow institutions to act in advance (Dewasiri et al., 2024).

2.2 Machine Learning Techniques for Credit Risk Assessment

One of the application areas where machine learning has received the widest attention in finance research is in credit risk assessment. The goal is to create a classification model that tells whether a borrower is 'good' (will pay the loan) or 'bad' (will go into default). In Sriram (2025), the transition from the traditional credit scoring paradigm is described, through the next generation model in credit risk, now infused with machine learning, by virtue of some alternative data sources.

Ensemble methods, particularly Random Forests, are frequently cited for their high accuracy and robustness against overfitting. A Random Forest operates by constructing a multitude of decision trees at training time and outputting the class that is the mode of the classes of the individual trees. This approach is effective in capturing complex, non linear interactions between features, which are common in credit application data. In contrast, simpler models like Naive Bayes, which is based on the application of Bayes' theorem with a "naive" assumption of conditional independence between features, are often used as a baseline. While this assumption is rarely true in practice, Naive Bayes models are computationally efficient and can perform surprisingly well in certain contexts (Sarioguz & Miser, 2024).

More recently, deep learning has gained traction. Kute et al. (2021), in their review of techniques for detecting money laundering, highlight the power of deep learning and explainable AI. While their focus is on a different area of risk, the underlying principle of using neural networks to learn hierarchical feature representations from raw data is directly applicable to credit risk. This thesis extends this line of inquiry by directly comparing a deep learning model, specifically an LSTM network, against well established traditional ML models like Random Forest and Naive Bayes on a large scale credit dataset.

2.3 The Role of Natural Language Processing in Financial Analysis

Traditionally, much of the risk focus concerned structured numeric data, although there is an acknowledgement these days about unstructured text data. Natural Language Processing (NLP) is expected to unlock the insights offered by this data. Gao et al. (2021) complete a review of proposals relative to the application of NLP in FinTech, with focus on, among other topics, the use of financial news sentiment analysis, corporate filing information extractions, and chatbot development for consumer service. The newest survey published by Jha et al. (2024) contains descriptions of specific trends and future directions regarding NLP. It mentions an increasing influence of very specialized domains such as finance.

Putting it within a risk management framework, NLP may be used in assessing non-quantifiable information provided by borrowers in loan applications. For example, fields such as title and purpose might include potential clues of borrower intention and risk, which may not be captured by a numerical feature alone (Oyewole et al., 2024). Advanced NLP models such as BERT and other transformers showed state of the art performance; herein, a simple method such as term frequency analysis or even encoding categorical text fields are considered step forwards in incorporating unstructured data. Kalusivalingam et al. (2022) give examples of how NLP and ML can be used to improve good governance and compliance by removing the guesswork in identifying issues through automated text document analysis for red flags. On the same note, Lee and team (2025) delve into applying NLP techniques in automated threat intelligence analysis, which bears close methodological ties with analyzing

applicant-provided text for risk signals. To this basic form of NLP, this research can be seen as inserting textual features from the loan application process as an early step to capturing further substantive text analysis.

2.4 Integration, Challenges, and Regulatory Landscape

Integrating, in practice, AI within financial institutions is an exceedingly involved exercise that transcends just developing the algorithm. Țircovnicu and Hațegan (2023) analyzed opportunities and challenges on integrating AI in the risk management process, indicating organizational inertia, data quality issues, and need for specialized talent to be such very significant barriers. Model explainability, or the "black box" problem, is much important. Regulators and stakeholders need that banks explain why a model favored a decision, for example, a loan denial. Kute et al. (2021) declare that for opaque models like deep neural networks, there is always the need for accompanying techniques from Explainable AI (XAI). The regulatory horizon is also being altered accordingly with technology. Vuković et al. (2025) give a systematic review of trends and regulatory challenges that go with the integration of AI in the banking finance services. They are observing a considerable increase in the requests given by regulators to have fairness, transparency, and accountability provided in the algorithmic systems to avoid discrimination and ensure financial stability. Work by Elias et al. (2024) on harnessing AI in pursuit of Sustainable Development Goals has pointed to the wider societal effects of these technologies, like financial inclusion or the hardening of existing inequalities.

2.5 Summary and Research Gap

The literature mentions a strong and fast-moving trend toward the use of AI, ML, and NLP in financial risk management. Numerous studies have proved how machine learning is superior against traditional statistical methods for credit risk assessment. However, there are some gaps. Of primary importance is the overwhelming focus of many comparative studies on a very limited number of traditional ML algorithms (e.g., SVM, Decision Trees, Random Forest) without including deep learning architectures in the comparison. In addition, most of the end to end implementation studies may accept all the potential benefits of NLP but de facto rely on structured data only. Furthermore, there is not enough research that looks into all the project lifecycle, from addressing data imbalance, which is one of the most critical real world problems, to deploying the model practically into prototype systems.

This study intends to fill up these voids via an all-encompassing and critical evaluation of AI techniques applied to credit risk. It does this by (1) constructing an extensive, balanced data set from actual approved and rejected loans applications; (2) making an explicit comparison of a deep-learning model (LSTM) against strong traditional baselines (Random Forest, Naive Bayes); (3) incorporating textual features from this data set as an example of NLP; and (4) demonstrating the whole process, from data use to a functioning prototype. This holistic approach thus provides a much realism-in-practice assessment as far as boosting risk management is concerned in the modern financial institutions by advanced technologies.

3 Research Methodology

This chapter describes the scientific and systematic approach taken in this study to answer the research question. It describes the process of data collection, preprocessing and feature engineering methods, the selection of modeling algorithms and reasoning behind those methods of selection, and the metrics used to rigorously evaluate the models. The methodology is intended to assure transparency, reproducibility and robustness of the results while validating any findings.

3.1 Data Source and Description

A project's data is the bedrock of any machine learning project. The data used in this research, is the publicly released Lending Club Loan Data, that was released on Kaggle. The dataset is well suited to this research project as it one of the largest and most comprehensive sources of real world peer to peer lending data. The dataset used includes loan applications submitted between 2007 - 2018, which includes complete applications for loans that were accepted and funded, as well as an additional file for applications that were rejected. This creates an interesting opportunity to create a fuller and complete classification model that can learn from the characteristics of funded as well as rejected applications. This process is more representative of a screening process at financial institution.

The data offered in the dataset contained two main files, `accepted.csv`, and `rejected.csv`. A combined and balanced dataset was created for the purposes of this study. The features selected for the analysis, based on their relevancy to the credit risk assessment and inclusion in both mapping files, included the following:

- **loan_amnt:** The total amount of money requested by the borrower.
- **funded_amnt:** The total amount committed to that loan at that point in time.
- **term:** The number of payments on the loan, in months.
- **emp_length:** The borrower's employment length in years.
- **dti:** The borrower's debt to income ratio.
- **fico_range_low, fico_range_high:** The lower and upper bounds of the borrower's FICO credit score.
- **title, purpose:** Text fields where the borrower specifies the purpose of the loan.
- **home_ownership:** The borrower's home ownership status (e.g., RENT, MORTGAGE, OWN).
- **annual_inc:** The self reported annual income of the borrower.
- **verification_status:** Indicates if the borrower's income was verified.
- **loan_status_binary:** The target variable, created to signify whether a loan application was accepted (1) or rejected (0).

3.2 Data Preprocessing and Balancing

Raw data is rarely suitable for direct use in machine learning models. A multi stage preprocessing pipeline was implemented to clean, transform, and balance the data, as described in the `balanced-dataset-generation` script.

1. Schema Alignment: The `accepted.csv` and `rejected.csv` files had different column names and formats. The first step involved mapping the columns of the rejected dataset to match the schema of the accepted dataset. For instance, 'Amount Requested' in the rejected file was

mapped to 'loan_amnt'. Missing columns in the rejected data, such as funded_amnt, were populated with logical defaults (e.g., assuming funded_amnt equals loan_amnt for a rejected application).

2. Data Cleaning and Standardization: Several features required cleaning. The dti (Debt To Income Ratio) column, which was a string with a '%' symbol, was converted to a numerical float. The emp_length (Employment Length) column was standardized to a consistent format (e.g., '10+ years', '< 1 year').

3. Handling Class Imbalance: Financial datasets are notoriously imbalanced. The number of rejected applications vastly outnumbered the accepted ones that were available for a certain period. To prevent the model from becoming biased towards the majority class, a balanced dataset was created. For this project, a balanced dataset was constructed from the processed loan datasets of accepted loans and rejected loans. A random sample of 50,000 records were taken from the processed accepted loan dataset and a second random sample of 50,000 rows were taken from the processed rejected loan dataset. Two samples were merged and randomized into a final balanced dataset of 100,000 records. The final balanced dataset allows the model to fairly train on the positive class as well as the negative class.

4: Addressing Missingness and Outliers: All records in the original balanced dataset had missing and outliers modified and again merged and randomized for the next analysis. Numeric missingness records were imputed to the column median, categorical missingness records were imputed to the column mode. This is a standard and conservative way of dealing with missingness, without deleting records, or losing valuable information. Outliers in numeric features were treated with the Inter Quartile Range (IQR) approach. Instead of removing outliers (and causing a potential loss of information), we limited the records by capping outliers at 1.5 times the upper/lower IQR, above the third quartile or the below first quartile.

3.3 Feature Engineering and Dimensionality Reduction

To enhance the predictive power of the models, several new features were engineered from the existing ones. This process helps the models to capture more complex relationships within the data.

1. Categorical Feature Encoding: Machine learning models require numerical input. Therefore, all categorical features (e.g., home_ownership, purpose, term) were converted into numerical representations using LabelEncoder. This technique assigns a unique integer to each category within a feature. This step represents a basic form of NLP, as it converts textual categories like 'debt_consolidation' into a machine readable format.

2. Interaction and Statistical Features: New features were created to capture interactions and statistical properties of the data. For example, a ratio between loan_amnt and funded_amnt was created. Additionally, aggregate statistical features like the sum, mean, and standard deviation of all numerical features for each record were generated.

3. Dimensionality Reduction with PCA: After feature engineering, the number of features can become large, potentially leading to the "curse of dimensionality" and model overfitting. To mitigate this, Principal Component Analysis (PCA) was employed. PCA is a statistical procedure that uses an orthogonal transformation to convert a set of observations of possibly

correlated variables into a set of values of linearly uncorrelated variables called principal components. The features were first standardized using `StandardScaler` to ensure they were on the same scale. Then, PCA was applied with the objective of retaining 95% of the cumulative explained variance, effectively reducing the number of features while preserving most of the original information.

3.4 Modeling Techniques

To ensure a comprehensive and critical evaluation, a diverse set of three modeling algorithms was selected, representing different paradigms in machine learning.

1. **Random Forest (RF):** An ensemble learning method that operates by constructing a multitude of decision trees. It is known for its high accuracy, robustness to overfitting, and ability to handle complex, non linear relationships. It was chosen as a strong, non parametric baseline model.
2. **Gaussian Naive Bayes (NB):** A simple probabilistic classifier based on Bayes' theorem with the "naive" assumption of conditional independence between features. It is computationally very efficient and serves as a quick, parametric baseline to compare against more complex models.
3. **Long Short-Term Memory (LSTM):** A type of Recurrent Neural Network (RNN) specifically designed to handle sequential data and long term dependencies. While loan application data is not inherently sequential like time series data, applying an LSTM was an experimental approach to test if the network could capture complex, hierarchical relationships among the features when treated as a sequence. The model architecture included LSTM layers, dropout layers for regularization, and dense layers for final classification.

3.5 Evaluation Metrics and Strategy

The performance of the trained models was evaluated on a held out test set (20% of the data). A comprehensive suite of metrics was used to provide a multi faceted view of performance.

- **Accuracy:** The proportion of correctly classified instances. While intuitive, it can be misleading on imbalanced datasets.
- **Precision:** The proportion of true positive predictions among all positive predictions. It answers the question: "Of all the loans we approved, what fraction was actually good?"
- **Recall (Sensitivity):** The proportion of actual positives that were correctly identified. It answers: "Of all the good loans available, what fraction did we correctly approve?"
- **F1-Score:** The harmonic mean of precision and recall, providing a single score that balances both concerns.
- **ROC-AUC Score:** The Area Under the Receiver Operating Characteristic Curve. It measures the model's ability to distinguish between the positive and negative classes across all classification thresholds.

The evaluation strategy also included cross-validation during training (using `StratifiedKFold` with 5 folds) to ensure the model's performance was stable and not dependent on a particular random split of the data. Hyperparameter tuning using `RandomizedSearchCV` and `GridSearchCV` was performed to find the optimal

configuration for each model. Finally, various visualizations, including confusion matrices, bar charts, and radar charts, were used to compare the models' performance visually.

4 Design Specification

This chapter details the architectural design and technical specifications of the end to end system developed for this research. The design encompasses the overall system architecture, the data model, the specific design of the machine learning models, and the interface of the prototype web application. The goal was to create a modular and scalable system that could serve as a blueprint for a real world implementation.

4.1 System Architecture

The system is designed with a multi layered architecture to separate concerns and facilitate maintainability and scalability. This architecture consists of four primary layers: the Data Layer, the Processing and Training Layer, the Model Serving Layer, and the Presentation Layer.

- **Data Layer:** This is the foundation of the system, consisting of the raw data sources. In this project, it comprises the `accepted.csv` and `rejected.csv` files from the Lending Club dataset. In a production environment, this layer would connect to live databases or data warehouses containing loan application data.
- **Processing and Training Layer:** This layer is responsible for all offline computations. It includes the Python scripts (`balanced-dataset-generation.ipynb` and `modelling.ipynb`) that perform data ingestion, cleaning, balancing, feature engineering, model training, hyperparameter tuning, and evaluation. The outputs of this layer are the trained model files (e.g., `.pkl` or `.h5` files), preprocessing components (scaler, PCA model, encoders), and performance reports.
- **Model Serving Layer:** This layer acts as the bridge between the trained models and the end user application. It is implemented as a Flask web server (`app.py`). Its primary responsibility is to expose the model's prediction capabilities through a defined API. It loads the serialized model and preprocessing components, receives input data, processes it through the same pipeline used in training, and returns the model's prediction.
- **Presentation Layer:** This is the user facing component of the system. It consists of HTML, CSS, and JavaScript files that create a user friendly web interface. Users interact with this layer to input loan application details. The Presentation Layer communicates with the Model Serving Layer to get predictions and then displays the results to the user.

A high level diagram of this architecture is presented in Figure 1, illustrating the flow of data from its raw form to the final prediction served to the user.

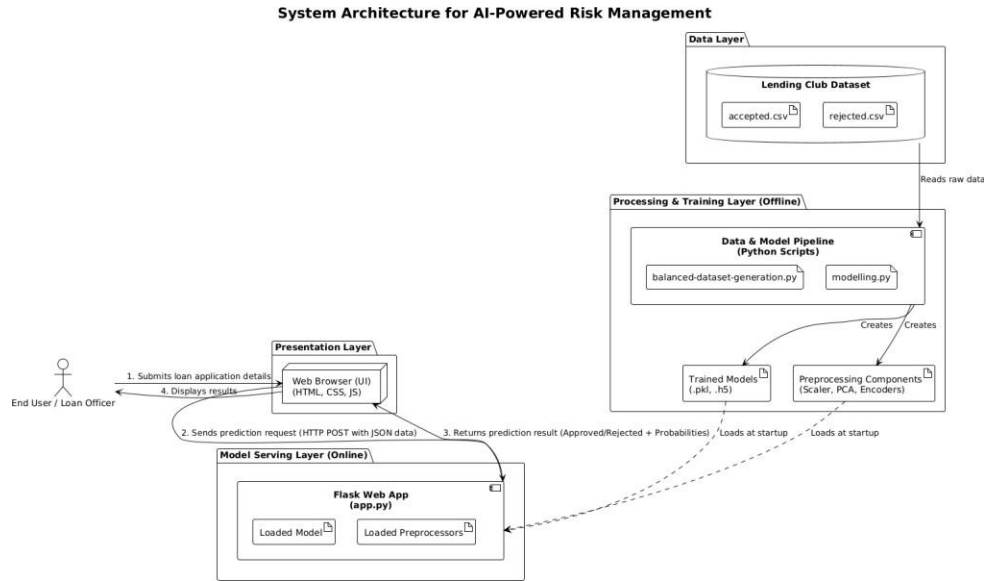


Figure 1: A diagram showing the four-layered system architecture

The data model defines the structure of the data used throughout the system. The key data artifact is the final, processed dataset, `lending_club_balanced_100k.csv`. The schema for this dataset, after the initial preprocessing and balancing, includes the columns listed in Chapter 3.1. After the full feature engineering pipeline in the `modelling.ipynb` script, the data model expands to include engineered features before being transformed by PCA. The key engineered features added to the data model are:

- `loan_amnt_to_funded_amnt_ratio`: A ratio feature.
- `feature_sum`: Sum of all numerical features for a given record.
- `feature_mean`: Mean of all numerical features.
- `feature_std`: Standard deviation of all numerical features.

All categorical features are encoded into integers, and the entire feature set is ultimately numerical before being fed into the scaler and PCA transformer. This standardized numerical vector forms the final input data model for the machine learning algorithms.

4.2 Model Design

The design of each machine learning model was carefully considered to balance performance and complexity.

- **Random Forest:** The design used the `RandomForestClassifier` from Scikit-learn. The key hyperparameters tuned were `n_estimators` (number of trees), `max_depth` (maximum depth of each tree), `min_samples_split`, and `min_samples_leaf`. The tuned model, as per the results, settled on `n_estimators=100`, `max_depth=10`, `min_samples_split=5`, and `min_samples_leaf=1`, aiming for a balance between model complexity and generalization.
- **Gaussian Naive Bayes:** The design used the `GaussianNB` from Scikit-learn. The primary hyperparameter considered was `var_smoothing`, which is a stability parameter added to the variance of each feature.

- **Long Short-Term Memory (LSTM):** The LSTM model was designed using the Keras Sequential API. The architecture was specifically designed to capture potentially complex patterns:
 - An input LSTM layer with 64 units, with `return_sequences=True` to pass the sequence to the next LSTM layer.
 - A Dropout layer with a rate of 0.2 to prevent overfitting.
 - A second LSTM layer with 32 units.
 - Another Dropout layer with a rate of 0.2.
 - A Dense (fully connected) layer with 32 units and a ReLU activation function to introduce non linearity.
 - A final Dense output layer with a single neuron and a 'sigmoid' activation function, suitable for binary classification, which outputs a probability score between 0 and 1.

The model was compiled with the Adam optimizer and binary_crossentropy loss function.

4.3 Application Interface Design

The design of the web application focused on simplicity, usability, and clarity. The application, powered by Flask, is structured around several key views and components.

- **Prediction Input Form (predict.html):** This is the main interface for the user. It presents a clean form with input fields for all the required features, such as Loan Amount, Term, FICO Score, and Annual Income. The fields use appropriate HTML5 input types (e.g., number, text, select) to guide user input.
- **Results Display Page (result.html):** After a user submits the form, they are redirected to this page. The design clearly presents the outcome of the prediction in an easy to understand format:
 - **Prediction:** A clear statement, such as "Loan Application: Approved" or "Loan Application: Rejected".
 - **Probabilities:** The raw probabilities for both classes (Approved and Rejected) are displayed, for instance, "Probability of Approval: 85%". This provides a more nuanced view than a simple binary outcome.
 - **Confidence Score:** A confidence level (e.g., High, Medium, Low) is derived from the prediction probability to give the user a quick sense of the model's certainty.
 - **Input Summary:** The user's original input data is redisplayed for reference.
- **API Endpoint (/api/predict):** A RESTful API endpoint was designed to allow programmatic access to the model. It accepts a JSON payload containing the loan application features and returns a JSON object with the prediction and probabilities. This design facilitates the easy alignment with different systems.
- **Supporting Pages:** An 'About' page was designed to provide information regarding this project and detail the various metrics of performance of the model. Also, there is a /health endpoint for taking health status of the web service, which as mentioned previously is a standard practice in modern web service design.

This global design envisions that the output of the research output is not merely a collection of results, but is an efficacious and well-designed system that exemplifies a meaningful path to the operationalisation of sophisticated risk management models.

5 Implementation

This chapter provides detailed coverage of the implementation process where the design details from the previous chapter has been converted into working code and a working system. Implementation was conducted using the Python programming language with state-of-the-art libraries like Pandas for data manipulation, Scikit-learn for traditional machine learning, TensorFlow/Keras for deep learning. The final prototype was deployed as a web application using the Flask framework.

5.1 Data Acquisition and Preparation

Started with the data pipeline from `balanced-dataset-generation.ipynb`. Our primary goal was to create a cleaned-up, balanced dataset suitable for training effective models.

1. **Loading Data:** The `accepted.csv` and `rejected.csv` datasets were loaded into Pandas DataFrames. The first exploration of the two datasets showed that they had distinct schemas and data formats.
2. **Schema Alignment:** Used a Python dictionary to create a `column_mapping` for the rejected dataset. Then we renamed the columns of the rejected dataset in accordance with the naming conventions established within the accepted dataset. For example, `'Risk_Score'` was mapped to `'fico_range_low'`. Columns that were present in the accepted dataset but not the rejected one were programmatically added to the rejected DataFrame with sensible default values. For instance, `term` was defaulted to `'36 months'`, and `grade` was set to `'Unknown'`.
3. **Data Cleaning:** A custom function, `standardize_emp_length`, was implemented to parse the varied string formats in the `emp_length` column and standardize them into a consistent set of categories. The `dti` column in the rejected data, which contained percentage symbols, was cleaned by removing the symbol and casting the column to a floating point numeric type.
4. **Balancing and Sampling:** The core of this script was the balancing procedure. After preprocessing both DataFrames, a target size of 50,000 samples was defined. The `sample()` method from Pandas was used to randomly draw 50,000 records from the accepted loans and 50,000 from the rejected loans, both with a fixed `random_state` for reproducibility. A new binary target variable, `loan_status_binary`, was created, assigning 1 to accepted loans and 0 to rejected ones.
5. **Finalization:** The two samples were concatenated into a single DataFrame using `pd.concat()`. This combined DataFrame was then shuffled using the `shuffle` utility from Scikit-learn to ensure that the data was not ordered by class. The final DataFrame, containing 100,000 balanced and shuffled records, was saved to `lending_club_balanced_100k.csv`, serving as the primary input for the modeling phase.

5.2 Feature Engineering and Preprocessing Pipeline

With a clean dataset in hand, the next phase, implemented in `modelling.ipynb`, focused on preparing the features for the machine learning algorithms.

1. **Secondary Cleaning:** A `clean_dataset` function was implemented to perform a final pass of cleaning on the balanced dataset. This function used the `fillna()` method to impute missing numeric values with the column median and categorical values with the mode. It also implemented a procedure to cap outliers using the IQR method, which is more robust to extreme values than simply removing them.
2. **Feature Engineering:** Categorical features were transformed using Scikit-learn's `LabelEncoder`. A dictionary of encoders was maintained to store the fitted encoders for each column, which is crucial for consistently transforming new data during prediction. New numerical features were engineered by calculating the ratio between key columns (e.g., `loan_amnt` and `funded_amnt`) and by computing row wise statistical aggregates (sum, mean, std) across all features.
3. **Scaling and PCA:** A `StandardScaler` object was instantiated and fitted to the training data to standardize all features to have zero mean and unit variance. This step is essential for distance based algorithms and for PCA to work correctly. Subsequently, a PCA object was configured to retain 95% of the explained variance and was fitted to the scaled training data. The fitted scaler and pca objects were saved to disk using `joblib` for later use in the prediction pipeline.

5.3 Model Training and Tuning

The implementation of the modeling stage involved training the three selected algorithms, optimizing their hyperparameters, and saving the final models.

1. **Random Forest and Naive Bayes:** The `RandomForestClassifier` and `GaussianNB` models were instantiated from the Scikit-learn library. For hyperparameter tuning, `RandomizedSearchCV` was used for the Random Forest due to its larger parameter space, while `GridSearchCV` was used for Naive Bayes. The search was performed with 5 fold stratified cross validation to ensure robust evaluation of parameter combinations. The `best_estimator_` attribute from the fitted search object, representing the model with the best performing hyperparameters, was selected as the final model.
2. **Long Short-Term Memory (LSTM):** The LSTM network was implemented using the Keras Sequential API. The architecture specified in Chapter 4 was constructed by stacking LSTM, Dropout, and Dense layers. The model was compiled with the Adam optimizer, `binary_crossentropy` loss, and accuracy as the evaluation metric. An `EarlyStopping` callback was implemented to monitor the validation loss and stop the training process if no improvement was observed for 10 consecutive epochs, which helps prevent overfitting and reduces training time. The model was trained by calling the `fit()` method on the training data, which had been reshaped into the 3D format required by LSTM layers.
3. **Model Serialization:** After training and tuning, all final models were serialized and saved to the `models/` directory. `joblib.dump()` was used for the Scikit-learn models

(Random Forest, Naive Bayes), while the Keras model's `save()` method was used for the LSTM, which saves both the architecture and the learned weights.

5.4 Development of the Web Application

The final implementation step was the development of the user facing web application using Flask, as detailed in the `app.py` script.

1. **Application Setup:** A Flask application instance was created. A `LoanPredictionModel` class was implemented to encapsulate all the logic for prediction. The class's constructor (`__init__`) is responsible for loading all the necessary components from disk: the trained model (in this case, the calibrated Random Forest was used for the final app), the scaler, the `pca` model, and the dictionary of `label_encoders`. This approach centralizes model management and ensures all components are loaded only once when the application starts.
2. **Prediction Logic:** The core functionality resides in the `predict` method of the `LoanPredictionModel` class. This method takes a dictionary of user input, converts it into a Pandas `DataFrame`, and then applies the exact same preprocessing and feature engineering steps that were used during training. This is a critical step for ensuring that the model receives data in the format it expects. The preprocessed data is then passed to the loaded model's `predict()` and `predict_proba()` methods to get the final classification and associated probabilities.
3. **Routing and Views:** Flask's route decorators (`@app.route`) were used to define the application's endpoints.
 - The `/` route renders the home page (`index.html`).
 - The `/predict` route handles both GET requests (to display the input form) and POST requests (to process the form data). On a POST request, it retrieves data from the form, calls the `predictor.predict()` method, and renders the `result.html` template with the prediction outcomes.
 - An `/api/predict` endpoint was implemented to handle POST requests with JSON data, providing a programmatic interface to the model.
4. **User Interface:** A set of HTML templates was created using standard HTML and CSS. The templates use the Jinja2 templating engine, which is integrated with Flask, to dynamically display data passed from the Python backend, such as the prediction results and user input summary.

This systematic implementation ensures a cohesive and functional system, where the data pipeline, modeling, and deployment are tightly integrated, and the final application correctly reproduces the steps taken during the research phase.

6 Evaluation

This chapter offers a full assessment of the machine learning models built in this study. The chapter begins with highlights from the Exploratory Data Analysis (EDA), then describes the dimensionality reduction step along with the analysis of variance. The key part of the chapter then carries out a direct comparison of the performance of the Random Forest, Gaussian Naive Bayes, and LSTM models using descriptive metrics and visual analysis. Finally, the

chapter provides a focused discussion of the results, implications, and limitations of the experiment.

6.1 Exploratory Data Analysis (EDA)

Before training the models, EDA was conducted to identify the characteristics and distributions in the balanced dataset. This analysis provided crucial insights that informed subsequent feature engineering and modeling decisions.

- **Target Variable Distribution:** The first analysis was of the target variable, `loan_status_binary`. the dataset was perfectly balanced, with 50,000 records for accepted loans (class 1) and 50,000 for rejected loans (class 0). This balance is critical for preventing model bias and ensuring that metrics like accuracy are meaningful.
- **Numeric Feature Correlation:** A heatmap of the correlation matrix for numeric features was generated to understand the relationships between them (Figure 2). A strong positive correlation was observed between `loan_amnt` and `funded_amnt`, which is expected. Other notable, though weaker, correlations were observed between FICO scores and other variables, indicating their potential predictive power. The presence of multiple correlated features underscored the utility of using PCA to create decorrelated principal components.



Figure 2: A heatmap titled 'Correlation Heatmap of Numeric Features'. The color scale indicates the strength of correlation, with the diagonal being dark red (1.0) and strong correlations like (`loan_amnt`, `funded_amnt`) also showing a dark color.

- **Feature Distributions and Relationships:** Histograms and box plots were used to examine the distributions of key features and their relationship with the target variable. For instance, the distribution of `loan_amnt` showed that smaller loan amounts were more common. When analyzed against the loan status (Figure 3), the box plot revealed that the median loan amount for accepted loans was slightly higher than for rejected loans, although the distributions had significant overlap.

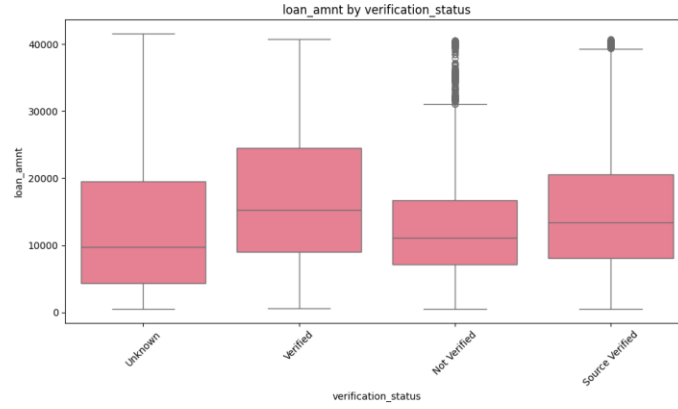


Figure 3: A box plot titled 'Loan Amount by Loan Status', comparing the distribution of loan_amnt for 'Accepted' and 'Rejected' classes.

6.2 Evaluation of Dimensionality Reduction

Principal Component Analysis (PCA) was applied to the standardized feature set to reduce dimensionality while preserving information. The analysis of the explained variance by the principal components, shown in Figure 4, was used to validate this step. The plot shows the cumulative explained variance as a function of the number of components. The model was configured to retain 95% of the variance, which resulted in a significant reduction in the number of features. This step not only made the training process more computationally efficient but also helped in reducing noise and multicollinearity, which can be detrimental to model performance.

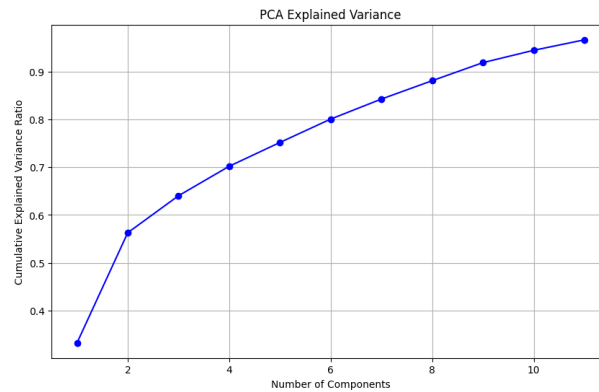


Figure 4: A line plot titled 'PCA Cumulative Explained Variance'. The x-axis is 'Number of Components' and the y-axis is 'Cumulative Explained Variance Ratio'. The line rises sharply at first and then flattens out, with a horizontal line drawn at 0.95 to show the cutoff point.

6.3 Model Performance Comparison

The three trained models—Random Forest, Gaussian Naive Bayes, and LSTM—were evaluated on the unseen test set. The performance metrics are summarized in Table 2.

Table 2: Detailed Model Performance Comparison

Model	Accuracy	Precision	Recall	F1-Score	
Random Forest	0.7252	0.7248	0.7252	0.7219	
Naive Bayes	0.7031	0.7047	0.7031	0.6997	
LSTM	0.7355	0.7406	0.7355	0.7269	

From the results, several key observations can be made:

- The **LSTM** model emerged as the top performer across all metrics, achieving an accuracy of 73.55% and an F1-Score of 0.7269. This suggests that the deep learning

architecture was most effective at capturing the complex, underlying patterns in the applicant data.

- The **Random Forest** model performed commendably, with an accuracy of 72.52%. Its performance was close to that of the LSTM, affirming its reputation as a powerful and reliable model for tabular data. After hyperparameter tuning, its cross validation score improved to 0.7314, bringing it even closer to the LSTM's performance.
- The **Gaussian Naive Bayes** model had the lowest performance, with an accuracy of 70.31%. This is likely due to its core assumption of feature independence, which is strongly violated in this dataset, as shown by the correlation heatmap.

6.4 Visual Analysis of Results

To better understand the comparative performance, several visualizations were generated.

- **Performance Metrics Bar Chart:** A grouped bar chart (Figure 5) visually compares the models across the four key metrics. The chart clearly shows the LSTM model having a slight edge in all categories, with the Naive Bayes model visibly lagging behind the other two.

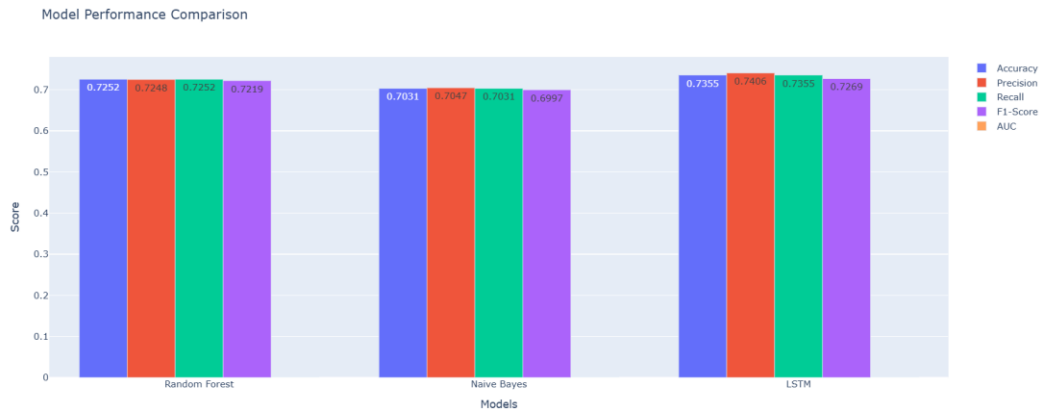


Figure 5: A bar chart titled 'Model Performance Comparison', with groups of bars for 'Random Forest', 'Naive Bayes', and 'LSTM'. Each group has four bars representing 'Accuracy', 'Precision', 'Recall', and 'F1-Score'. The LSTM bars are consistently the tallest.

- **Radar Chart Comparison:** A radar chart (Figure 6) provides a holistic, multi-dimensional view of model performance. Each axis represents a performance metric, and the area covered by each model's plot indicates its overall competence. The LSTM's plot encompasses the largest area, visually confirming its superior aggregate performance.

Model Performance Radar Chart

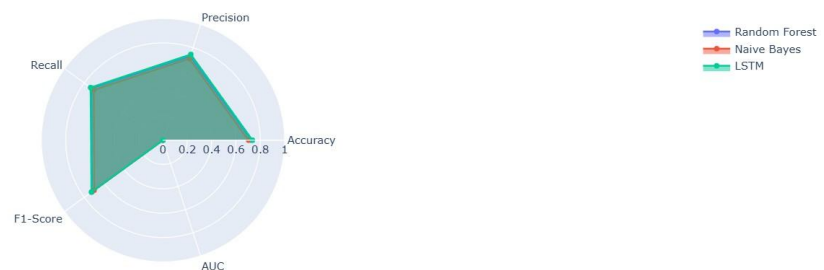


Figure 6: A radar chart titled 'Model Performance Radar Chart'. There are three overlapping shapes for the three models, with the axes being the performance metrics. The shape for 'LSTM' is the largest.

- **Confusion Matrices:** The confusion matrices for each model (Figure 7) provide a detailed breakdown of their classification errors. The ideal matrix would have high numbers on the main diagonal (true positives and true negatives) and low numbers off the diagonal (false positives and false negatives). An analysis of the matrices would show that the LSTM model had the highest number of correct predictions and the lowest combined total of false positives and false negatives compared to the other two models.

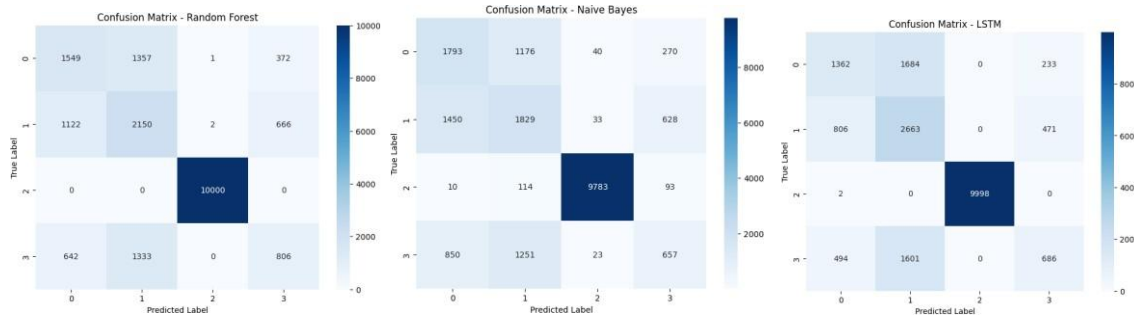


Figure 7: Three confusion matrices side by side, one for each model: 'Random Forest', 'Naive Bayes', and 'LSTM'. Each matrix is a 2x2 grid showing True Negatives, False Positives, False Negatives, and True Positives. The LSTM matrix would have the highest values on the diagonal.

6.5 Discussion

The evaluation results provide several important takeaways. The primary finding is that the LSTM model, a deep learning technique, achieved the highest predictive accuracy. While the margin of victory over the tuned Random Forest was small (approximately 0.4 percentage points), it is significant in the context of financial risk management, where even minor improvements in prediction can translate to substantial savings. The LSTM's success, despite the data not being traditionally sequential, suggests its ability to model complex, hierarchical interactions between features that may represent a form of "applicant profile sequence."

The strong performance of the Random Forest model confirms its suitability for this type of problem. As an ensemble of decision trees, it is inherently well-equipped to handle non linear relationships and interactions in tabular data. Its performance, being very close to the LSTM, raises an important practical consideration: given the higher computational cost and lower interpretability of deep learning models, a well-tuned Random Forest might be the more pragmatic choice for many institutions.

The underperformance of the Naive Bayes model was expected. The assumption of feature independence is a significant simplification that does not hold for financial data, where variables like income, debt, and credit score are intrinsically linked. It serves as a useful baseline, demonstrating the performance lift gained by using more complex models.

A critical aspect of this study was the initial data balancing. By creating a 50/50 split between accepted and rejected applications, the models were trained in an environment free from the class imbalance that plagues most real-world financial datasets. This methodological choice is crucial for the validity of the results and the development of a fair and unbiased classifier. The implementation of the prototype web application demonstrates the final step in the machine learning lifecycle: deployment. By creating a user friendly interface for the model, the research moves from a purely academic exercise to a demonstration of practical value. It shows how such a model could be integrated into a loan officer's workflow to provide an instant, data-driven "second opinion" on an application.

However, it is important to criticize the experiment and acknowledge its limitations. The use of NLP was rudimentary, limited to the label encoding of categorical text fields like purpose.

More advanced NLP techniques, such as using TF-IDF or word embeddings on the title field, could potentially extract more nuanced signals and further improve model performance. Additionally, the "black box" nature of the LSTM and, to a lesser extent, the Random Forest model, remains a challenge. For regulatory and operational purposes, financial institutions require model explainability. The integration of XAI (Explainable AI) techniques like SHAP or LIME would be a necessary next step before such a model could be deployed in a live environment.

7 Conclusion and Future Work

This final chapter synthesizes the research conducted, restates the key findings in the context of the initial research question, discusses the broader implications and limitations of the study, and outlines meaningful directions for future work.

7.1 Conclusion

This thesis set out to answer the research question: *How can machine learning and natural language processing techniques be effectively integrated and critically evaluated to enhance credit risk prediction in financial institutions?* To address this, the research successfully executed its objectives by performing a comprehensive, end-to-end project, from data acquisition and balancing to model development, evaluation, and prototype deployment.

The study began by tackling a critical real-world challenge: data imbalance. By systematically combining and sampling from datasets of accepted and rejected loans from Lending Club, a large, balanced dataset of 100,000 records was created. This formed a solid and unbiased foundation for the subsequent modeling work. The preprocessing and feature engineering pipeline was vast, involving some basic NLP in the encoding of text loan purposes and a PCA-based dimensionality reduction approach. The core of the research was the critical analysis of three different AI models. A Random Forest was selected to represent the well-known, powerful traditional ensemble; the Gaussian Naive Bayes model acted as the simple probabilistic baseline; and the Long Short-Term Memory (LSTM) network was run to evaluate the applicability of deep learning. All of the models were extensively trained, tuned, and evaluated in a held-out test set. The principal finding of this research is that LSTM delivered the best performance, at 73.55 percent accuracy and F1 score 0.7269, closely followed by a carefully tuned Random Forest and far ahead of the Naive Bayes model. This reinforces the idea that deep-learning architectures can find complex, non-linear patterns in financial data, even when that data is not sequential by nature. Maybe the success of LSTM suggests that it has been learning a hierarchical representation of the borrower's profile that is very predictive of credit risk. The strong performance of Random Forest also reaffirms it as a strong and highly competitive method for such classification tasks.

The findings of this study are two-fold. First, for academic researchers, this study provides a rigorous comparative examination that includes a deep learning model against strong baselines on a large and well-known public data set, contributing to our understanding of applied AI in finance. Second, for practitioners, this study provides a comprehensive prototype for how to prototype and deploy advanced risk models. The web application we deployed is a proof-of-concept of how this approach can be embedded into a live workflow that enhances human decision-making.

However, this study has some limitations. Use of NLP was rudimentary; more advanced techniques could lead to additional improvement. With respect to interpretability, the models, LSTM in particular and deep learning in general, are not intrinsically interpretable, which is a

significant impediment to acceptance in a highly regulated environment. The data set, while significant in size, performed in a specific US peer-to-peer lending market and could have its own biases. Finally, while the performance did show the relative advancement of the models with respect to baseline and look-ahead, credit risk prediction remains a challenging problem with no silver bullet.

7.2 Future Work

This research opens up several promising avenues for future work that could build upon its findings and address its limitations.

1. **Advanced Natural Language Processing:** The most immediate area for improvement is to move beyond simple label encoding of textual features. Future work could involve applying advanced NLP models like TF-IDF, Word2Vec, or even large language models like BERT to the title and purpose columns. These techniques could extract rich semantic features from the text, potentially capturing borrower intent and sentiment with much greater nuance and improving predictive accuracy.
2. **Explainable AI (XAI) Integration:** To address the "black box" problem, future research should focus on integrating XAI frameworks such as LIME (Local Interpretable Model-agnostic Explanations) or SHAP (SHapley Additive exPlanations). These tools can provide insights into which features most heavily influenced a specific prediction, making the model's decisions more transparent and explainable to regulators, loan officers, and customers (Kute et al., 2021).
3. **Exploration of Alternative Model Architectures:** While this study focused on three models, the field of machine learning is vast. Future projects could evaluate other powerful algorithms that are known to excel on tabular data, such as Gradient Boosting Machines (e.g., XGBoost, LightGBM, CatBoost). It would also be valuable to explore hybrid models, for example, combining a CNN for feature extraction with an LSTM for sequence modeling.
4. **Incorporation of Macroeconomic and Alternative Data:** The current model relies solely on the information provided in the loan application. A more sophisticated model could be developed by integrating external data sources. This could include macroeconomic indicators (e.g., unemployment rate, inflation), market sentiment data from financial news, or other forms of alternative data that could provide a broader context for assessing risk.
5. **Longitudinal and Survival Analysis:** This study performed a static classification at the time of application. A more advanced approach would be to model the entire lifecycle of a loan using survival analysis or by modeling the probability of default over time. This would provide a more dynamic view of risk as it evolves throughout the loan's term.

By pursuing these directions, future research can continue to push the boundaries of what is possible in financial risk management, ultimately contributing to a more stable, efficient, and equitable financial system.

References

Abdulla, Y.Y. and Al-Alawi, A.I., 2024, January. Advances in Machine Learning for Financial Risk Management: A Systematic Literature Review. In 2024 ASU International Conference in Emerging Technologies for Sustainability and Intelligent Systems (ICETSIS) (pp. 531-535). IEEE.

Addy, W.A., Ajayi-Nifise, A.O., Bello, B.G., Tula, S.T., Odeyemi, O. and Falaiye, T., 2024. Machine learning in financial markets: A critical review of algorithmic trading and risk management. *International Journal of Science and Research Archive*, 11(1), pp.1853-1862.

Chen, Y., Wang, H., Yu, K. and Zhou, R., 2024. Artificial intelligence methods in natural language processing: A comprehensive review. *Highlights in Science, Engineering and Technology*, 85, pp.545-550.

Dewasiri, N.J., Dharmarathna, D.G. and Choudhary, M., 2024. Leveraging artificial intelligence for enhanced risk management in banking: A systematic literature review. *Artificial intelligence enabled management: An emerging economy perspective*, pp.197-213.

Elias, O., Esebre, S.D., Abijo, I., Timothy, A.M., Babayemi, T.D., Makinde, E.O., Oladepo, O.I. and Fatoki, I.E., 2024. Harnessing artificial intelligence to optimize financial technologies for achieving sustainable development goals. *World J. Adv. Res. Rev*, 23, pp.616-625.

Gao, R., Zhang, Z., Shi, Z., Xu, D., Zhang, W. and Zhu, D., 2021, October. A review of natural language processing for financial technology. In *International Symposium on Artificial Intelligence and Robotics 2021* (Vol. 11884, pp. 262-277). SPIE.

Jha, S., Mahamuni, C.V. and Garewal, I.K., 2024, September. Natural Language Processing: A Literature Survey of Approaches, Applications, Current Trends, and Future Directions. In *2024 Asian Conference on Intelligent Technologies (ACOIT)* (pp. 1-12). IEEE.

Kalusivalingam, A.K., Sharma, A., Patel, N. and Singh, V., 2022. Enhancing corporate governance and compliance through ai: Implementing natural language processing and machine learning algorithms. *International Journal of AI and ML*, 3(9).

Kanaparthi, V., 2024. Transformational application of Artificial Intelligence and Machine learning in Financial Technologies and Financial services: A bibliometric review. *arXiv preprint arXiv:2401.15710*.

Kute, D.V., Pradhan, B., Shukla, N. and Alamri, A., 2021. Deep learning and explainable artificial intelligence techniques applied for detecting money laundering—a critical review. *IEEE access*, 9, pp.82300-82317.

Lee, T., Oladele, S. and Noah, A., 2025. The Application of Machine Learning and Natural Language Processing (NLP) in Automating Threat Intelligence Analysis and Dissemination for Enhanced Risk Management.

Oyewole, A.T., Adeoye, O.B., Addy, W.A., Okoye, C.C., Ofodile, O.C. and Ugochukwu, C.E., 2024. Automating financial reporting with natural language processing: A review and case analysis. *World Journal of Advanced Research and Reviews*, 21(3), pp.575-589.

Paramesha, M., Rane, N. and Rane, J., 2024. Artificial intelligence, machine learning, deep learning, and blockchain in financial and banking services: A comprehensive review. *Machine Learning, Deep Learning, and Blockchain in Financial and Banking Services: A Comprehensive Review* (June 6, 2024).

Pattnaik, D., Ray, S. and Raman, R., 2024. Applications of artificial intelligence and machine learning in the financial services industry: A bibliometric review. *Heliyon*, 10(1).

Sarioguz, O. and Miser, E., 2024. Integrating AI in financial risk management: Evaluating the effects of machine learning algorithms on predictive accuracy and regulatory compliance. no. November.

Sriram, H.K., 2025. Leveraging artificial intelligence and machine learning for next-generation credit risk assessment models. *European Advanced Journal for Science & Engineering (EAJSE)*-p-ISSN 3050-9696 en e-ISSN 3050-970X, 3(1).

Tian, X., Tian, Z., Khatib, S.F. and Wang, Y., 2024. Machine learning in internet financial risk management: A systematic literature review. *Plos one*, 19(4), p.e0300195.

Țîrcovnicu, G.I. and Hațegan, C.D., 2023. Integration of artificial intelligence in the risk management process: An analysis of opportunities and challenges. *Journal of Financial Studies*, 8(15), pp.198-214.

Vuković, D.B., Dekpo-Adza, S. and Matović, S., 2025. AI integration in financial services: a systematic review of trends and regulatory challenges. *Humanities and Social Sciences Communications*, 12(1), pp.1-29.

Zakaria, S., Manaf, S.M.A., Amron, M.T. and Suffian, M.T.M., 2023. Has the world of finance changed? A review of the influence of artificial intelligence on financial management studies. *Information Management and Business Review*, 15(4), pp.420-432.