



Predicting Financial Vulnerability in the Portuguese Population: A Machine Learning Approach

Guilherme Lopes de Andrade Gusmão

Dissertation written under the supervision of professor Domenico Fabrizi

Dissertation submitted in partial fulfilment of requirements for
the MSc in Business Analytics, at the Universidade Católica Portuguesa,
May 2026.

Abstract

Title: Predicting Financial Vulnerability in the Portuguese Population: A Machine Learning Approach

Author: Guilherme Lopes de Andrade Gusmão

Keywords: financial vulnerability; financial literacy; machine learning; Portugal; survey data; XGBoost; predictive modelling

Financial vulnerability is a real concern in Portugal with many households struggling to cope with a sudden shock equivalent to one month of income. This dissertation explores whether such vulnerability can be reliably predicted from survey responses, and which characteristics drive this prediction. Based on the 4th Survey on the Financial Literacy of the Portuguese Population (Banco de Portugal, CMVM and ASF, 2023), six supervised classifiers are trained on socio-demographic, behavioural, knowledge and attitudinal variables, with the binary target derived from a self-assessed item capturing the ability to absorb an expense equal to one month of personal income without borrowing. The performance of the models is assessed through nested cross-validation and an independent hold-out test. XGBoost emerges as the strongest model (test ROC-AUC of 0.913; PR-AUC of 0.954, computed with the resilient class as the positive label), narrowly ahead of gradient boosting and random forest. The most informative predictors are not knowledge items, as much of the literacy literature would predict, but rather concrete savings behaviour, retirement planning confidence and household income. An error analysis suggests that the population most likely to be misclassified as resilient consists of mid-income respondents who hold retirement products and intend to fund their retirement from personal savings, but are still unable to absorb a one-month shock, a profile worth targeting for financial education.

Resumo

Título: Previsão da Vulnerabilidade Financeira na População Portuguesa: Uma Abordagem Baseada em Aprendizagem Automática

Autor: Guilherme Lopes de Andrade Gusmão

Palavras-chave: vulnerabilidade financeira; literacia financeira; aprendizagem automática; Portugal; dados de inquérito; XGBoost; modelação preditiva

A vulnerabilidade financeira é uma preocupação real em Portugal, com muitos agregados familiares a lutarem para lidar com um choque inesperado equivalente a um mês de rendimento. Esta dissertação explora se essa vulnerabilidade pode ser prevista de forma fiável através de respostas a inquéritos e quais as características que impulsionam essa previsão. Com base no 4.º Inquérito à Literacia Financeira da População Portuguesa (Banco de Portugal, CMVM e ASF, 2023), treinaram-se seis classificadores supervisionados com variáveis sociodemográficas, comportamentais, de conhecimento e de atitude. O target binário deriva da capacidade autoavaliada de suportar uma despesa igual a um mês de rendimento sem recorrer a empréstimos. O desempenho dos modelos foi avaliado via validação cruzada aninhada e um teste independente (hold-out). O XGBoost destacou-se como o modelo mais forte (ROC-AUC de 0,913; PR-AUC de 0,954, com a classe resiliente como positiva), seguido de perto pelo gradient boosting e random forest. Os preditores mais informativos não são os itens de conhecimento, ao contrário do previsto pela literatura, mas sim o comportamento de poupança, a confiança no planeamento da reforma e o rendimento familiar. Uma análise de erros sugere que a população mais erradamente classificada como resiliente engloba respondentes de rendimento médio com produtos de reforma e intenção de se autofinanciarem, mas incapazes de absorver o choque, um perfil prioritário para a educação financeira.

Acknowledgments

With this dissertation, a long chapter of my life comes to an end. More than the two years of this master's, it marks the end of three years of undergraduate studies and the twelve years of school before them. Looking back, I am left with an enormous sense of gratitude toward the many people I met along the way, those who stayed, who I am lucky enough to call my friends today, and from whom I learned far more than any classroom ever taught me.

I am sincerely grateful to my supervisor, Professor Domenico Fabrizi, whose guidance, patience, and steady feedback shaped this work at every stage.

To my oldest friends, Miguel T., Miguel F., Miguel R., Afonso, and Helena, thank you for the years that reach back further than I can easily remember; you have been there through every version of me, and that constancy is something I will never take for granted. To António, Carolina, Maria, and Lourenço, with whom I had the privilege of working side by side at JDC, those years would not have been the same without you. A special word goes to André and Pedro, who shared a roof with me through these years and lived this journey alongside me day after day.

To Maria Vitória, who shared a large part of this journey with me, thank you for everything those years gave me and for all you taught me along the way. It was a chapter I'll always look back on with gratitude. And to her family, Tio Alexandre, Tia Cláudia, António, and Kiko, thank you for having welcomed me as one of your own from the very first day.

My deepest gratitude, though, belongs to my family. To my mother: you gave up more than I will ever know so that I could have everything. You always believed in me. Everything I am, I am because of you, and this is as much yours as it is mine. I love you. To my grandmother, Nã, whom I love so much, thank you for your constant presence, for your tenderness, and for being my safe harbour, no matter what. And to my father: I miss you every day, and that has never gotten easier. I wish more than anything that you could be here in person to see this, but I know you are watching over me from above, where God has you now. I carry you with me in everything I do, and I love you, always. This one is for you, Dad.

Contents

Abstract	1
Resumo	2
Acknowledgments	3
List of Abbreviations	9
1 Introduction	1
1.1 Research Context	1
1.2 Research Problem and Questions	1
1.3 Significance of the Study	2
1.4 Structure of the Dissertation	2
2 Literature Review	3
2.1 Financial Vulnerability: Definition, Measurement and Determinants	3
2.2 Financial Literacy and Its Relationship to Financial Vulnerability	4
2.2.1 The Three-Component Framework	4
2.2.2 Knowledge, Behaviour and Outcomes	4
2.2.3 The Portuguese Evidence	5
2.2.4 Financial Attitudes and the Behavioural Channel	5
2.3 Machine Learning Applied to Financial Behaviour Prediction	6
2.3.1 Prior Applications in Household Finance	6
2.3.2 Why Machine Learning May Add Value in This Context	6
2.3.3 Key Algorithms	7
2.4 Research Gaps and Positioning of This Study	7
3 Data and Methodology	9
3.1 Data Source	9
3.2 Target Variable	9
3.3 Feature Engineering	10
3.3.1 Knowledge Score Construction	10
3.3.2 Multicoded Variables	11
3.3.3 Ordinal Encoding	11
3.4 Data Cleaning	11
3.5 Modelling Pipeline	12
3.5.1 Preprocessing	12

3.5.2	Class Imbalance Handling	12
3.6	Model Selection and Evaluation	13
3.6.1	Nested Cross-Validation	13
3.6.2	Evaluation Metrics	13
4	Results	15
4.1	Exploratory Data Analysis	15
4.1.1	Target Distribution	15
4.1.2	Demographic and Economic Patterns	15
4.1.3	Knowledge Score Distribution	17
4.2	Model Performance Comparison	17
4.2.1	Nested Cross-Validation Results	17
4.2.2	Test Set Performance	18
4.2.3	ROC and Precision-Recall Curves	19
4.2.4	Overfitting Assessment	21
4.3	Feature Importance and Model Interpretability	22
4.3.1	Feature Importance: Random Forest	22
4.3.2	Logistic Regression Coefficients	23
4.3.3	Confusion Matrices	24
4.4	Profile of Vulnerable Respondents Classified as Resilient	25
4.5	Key Predictors of Financial Vulnerability	27
5	Discussion	29
5.1	Interpretation of Results in Light of the Literature	29
5.2	Policy Implications for Financial Education	31
5.3	Limitations	31
5.4	Future Research	33
6	Conclusion	34
	References	36
A	Variable Descriptions	38
B	Confusion Matrices	39
C	Hyperparameter Search Spaces	40

List of Figures

4.1	Distribution of the target variable (B6) in the final analytical sample after exclusions and listwise deletion of missing values. The resilient class accounts for approximately 69 per cent of observations.	15
4.2	Financial resilience (ability to cover an unexpected expense equivalent to one month's personal income) by gender, age band, education level and region of residence. Bars show the within-group share of resilient (Yes) and vulnerable (No) respondents in the analytical sample (n = 1,355). Percentages are row-wise within each category.	16
4.3	Distribution of the composite financial knowledge score (D2–D7) by target class. While resilient respondents score somewhat higher on average, the overlap between the two distributions is substantial.	17
4.4	Box plot of ROC-AUC scores across the five outer cross-validation folds for each model. XGBoost shows both the highest median and the tightest interquartile range.	18
4.5	Test-set ROC-AUC for all six models, presented as a bar chart. All ensemble methods exceed 0.900, with XGBoost narrowly leading. The chart complements the cross-validation distribution shown in Figure 4.4.	19
4.6	Receiver operating characteristic (ROC) curves for all six models on the held-out test set. The diagonal represents random guessing.	20
4.7	Precision-recall curves for all six models on the test set. Logistic regression and XGBoost achieve near-identical PR-AUC values despite differing in specificity and recall trade-offs.	20
4.8	Training versus test ROC-AUC gap for each model. The KNN gap reflects an artefact of distance-based weighting on the training set rather than genuine in-sample learning; among ensemble models, random forest generalises most cleanly.	21
4.9	Top 20 features by random forest impurity-based importance. Savings behaviour (B5), retirement confidence (B10) and income (D11) clearly outweigh objective knowledge scores.	23
4.10	Top 15 logistic regression coefficients sorted by absolute magnitude. Positive coefficients indicate a higher probability of being classified as resilient; negative coefficients indicate higher vulnerability.	24
4.11	Confusion matrices for all six models on the held-out test set (n = 271). Each cell reports the absolute count; row labels indicate the true class and column labels the predicted class.	25

4.12	Comparison of mean feature values for the 19 vulnerable respondents misclassified by XGBoost as resilient, the 65 correctly classified vulnerable respondents (TN) and the full test set, across selected numeric features.	27
B.1	Confusion matrices for all six models (test set, n = 271). XGBoost achieves the best balance between true positives (167) and true negatives (65), while logistic regression shows higher specificity at the cost of more false negatives.	39

List of Tables

4.1	Nested cross-validation results (outer 5-fold, inner RandomizedSearchCV with 30 iterations and 5-fold CV, scored on ROC-AUC). Models are sorted by mean ROC-AUC.	18
4.2	Model performance on the held-out test set (n = 271). The positive class is Resilient (1). Models are sorted by ROC-AUC.	19
4.3	Train versus test ROC-AUC gap as an indicator of overfitting. Smaller gaps indicate better generalisation.	21
4.4	Random forest feature importance (top 16 by impurity-based importance). Behavioural and attitudinal variables dominate the ranking; the composite knowledge score appears at position 16.	22
4.5	Profile of the 19 vulnerable respondents XGBoost classified as resilient, compared to the 65 correctly classified vulnerable respondents (TN) and the full test set. Numeric features report means; categorical features report modes.	26
A.1	Summary of variable groups used in the modelling pipeline, with survey section, key items and encoding approach.	38
C.1	Hyperparameter search spaces for each model (RandomizedSearchCV, n_iter=30, 5-fold stratified CV, scoring='roc_auc'). Models whose total grid size was smaller than 30 were searched exhaustively.	40

List of Abbreviations

ASF	Autoridade de Supervisão de Seguros e Fundos de Pensões
AUC	Area Under the Curve
BdP	Banco de Portugal
CMVM	Comissão do Mercado de Valores Mobiliários
CV	Cross-Validation
EDA	Exploratory Data Analysis
FN	False Negative
FP	False Positive
GBM	Gradient Boosting Machine
INE	Instituto Nacional de Estatística
INFE	International Network on Financial Education
kNN	k-Nearest Neighbours
LogReg	Logistic Regression
ML	Machine Learning
OECD	Organisation for Economic Co-operation and Development
PR-AUC	Precision–Recall Area Under the Curve
RF	Random Forest
ROC-AUC	Receiver Operating Characteristic Area Under the Curve
XGB	Extreme Gradient Boosting

1. Introduction

1.1 Research Context

The ability to resist a financial shock without going into debt is one of the most concrete markers of well-being. Successive waves of the Survey on the Financial Literacy of the Portuguese Population, conducted by Banco de Portugal (BdP), the Comissão do Mercado de Valores Mobiliários (CMVM) and the Autoridade de Supervisão de Seguros e Fundos de Pensões (ASF), have consistently shown that a significant portion of the Portuguese population would not be able to cover an unexpected expense equivalent to one month of personal income without external help (Conselho Nacional de Supervisores Financeiros, 2023a). The 2023 edition of the survey, conducted through 1,510 face-to-face interviews during January and February 2023, gives the latest cross-sectional picture of how Portuguese residents manage, plan and think about money.

Evidence from the United States and across the European Union suggests similar fragilities (Ampudia et al., 2016; Lusardi et al., 2011). What makes the case interesting for research is the simultaneous availability of granular data on financial knowledge, attitudes, behaviours and product holdings, all collected under a standardised OECD/INFE framework (OECD, 2022). That combination allows to move beyond describing who is vulnerable and investigate which dimensions of financial literacy, if any, are most closely associated with resilience. Existing analyses of the BdP survey rely on descriptive cross-tabulations or simple regression models, leaving open the question of whether modern classification techniques can extract richer predictive structure from the same data.

1.2 Research Problem and Questions

This dissertation attempts to answer the following question: *To what extent can financial vulnerability in the Portuguese population be predicted from survey-based measures of socio-demographic characteristics, financial behaviours, knowledge and attitudes, and which predictors matter most?*

The analysis is divided into three sub-questions. The first question asks which supervised learning algorithm, among six candidates, achieves the strongest performance when predicting the ability to withstand a one-month income shock. The second one concerns the nature of the predictors, whether are objective knowledge items, behavioural indicators or attitudinal self-assessments more informative? The third one explores the error structure of the best-performing model, asking who is most likely to be misclassified as financially resilient when they are, in fact, vulnerable.

1.3 Significance of the Study

Three considerations motivate this project. At a methodological level, the machine learning classifiers have proved to be valuable in related domains such as credit scoring and default prediction (Khandani et al., 2010; Moscatelli et al., 2020). The Portuguese case, however, has so far been studied almost exclusively through descriptive cross-tabulations and linear regression (Conselho Nacional de Supervisores Financeiros, 2023a), and no published analysis to the author’s knowledge applies a systematic comparison of supervised classifiers to the 2023 wave of the BdP/CMVM/ASF survey. This thesis tests whether that flexibility translates into measurable predictive gains.

At a substantive level, Portugal occupies a middle position in European comparisons of financial literacy (OECD, 2023). Understanding what distinguishes roughly one third of the survey respondents who report being unable to cover an unexpected expense (Conselho Nacional de Supervisores Financeiros, 2023a) from the remainder can give some insights into the intervention need in those cases.

At a practical level, a predictive model that performs well on unseen data could be integrated into screening tools used by social services, non-governmental organisations or supervisory authorities. Even a decision rule derived from model outputs may help to identify individuals who are at elevated risk of financial distress before a shock happens. This dissertation does not build such a tool, but it provides the empirical foundation on which one could be created.

1.4 Structure of the Dissertation

The remainder of this dissertation is organised as follows. Chapter 2 analyses the relevant literature across three domains: the conceptualisation and assessment of financial vulnerability, the relationship between financial literacy and resilience, and the application of machine learning to household finance. Chapter 3 addresses the data source, the construction of the target variable and the nested cross-validation framework used for model selection, while Chapter 4 reports the empirical findings, model comparisons, feature importance rankings and a profile of the misclassified vulnerable respondents. The interpretation of these findings, together with their policy implications for the Portuguese context and the study’s limitations, is the subject of Chapter 5. Chapter 6 closes the dissertation by going back to the research questions and pointing to directions for future work.

2. Literature Review

This study is based on three main areas of research. The first is related to the definition and measurement of financial vulnerability at the household level. The second addresses the relationship between financial literacy, understood through its tripartite structure of knowledge, behaviours and attitudes, and observed financial outcomes. The third looks at the growing use of machine learning in predicting financial behaviours and distress. Each of these strands is reviewed in turn below, followed by a discussion of the research gap that the present study aims to fill.

2.1 Financial Vulnerability: Definition, Measurement and Determinants

There is no universally accepted definition of financial vulnerability, and the concept has been operationalised in many different ways depending on the data and discipline involved. Some researchers focus on objective indicators such as net wealth, debt-to-income ratios or the presence of liquid savings buffers, while others privilege subjective measures like self-reported ability to cope with unexpected expenses or perceived financial stress (Anderloni et al., 2012). This difference is important because the two approaches do not always identify the same individuals as vulnerable. A household with modest income but no debt and a habit of precautionary savings may score well on a balance-sheet measure yet report high anxiety about money; conversely, a higher-income household carrying substantial mortgage debt may appear objectively leveraged while feeling more secure.

The measure used in this dissertation is more on the subjective end of the spectrum but has a concrete behavioural basis. In the question B6 of the BdP/CMVM/ASF survey, respondents are asked whether they would be able to personally face an unexpected expense equivalent to one month of their own income without borrowing or asking family and friends for help. It is also aligned with the OECD/INFE financial resilience framework, which treats the capacity to handle shocks as a core component of financial well-being (Conselho Nacional de Supervisores Financeiros, 2023a; OECD, 2022).

Research based on comparable questions in other countries has consistently identified a set of common correlates of financial fragility. Lusardi et al. (2011) documented that nearly half of American respondents reported they could not cover a hypothetical \$2,000 emergency expense, with vulnerability concentrated among lower-income, less-educated and younger households. In the European context, Ampudia et al. (2016) analysed data from the Household Finance and Consumption Survey to map fragility across euro-area countries, finding that it was shaped not only by income and wealth levels but also by the structure of household debt. More recently,

Brunetti et al. (2016) studied Italian households during the Great Recession and argued that vulnerability has a dynamic quality: shocks can push previously resilient families into fragile states, especially when precautionary savings are thin.

For Portugal specifically, the 2023 survey report documents that the financial resilience indicator improved relative to 2020, with the mean rising from 47.9 to 55.2 on a normalised 0–100 scale (Conselho Nacional de Supervisores Financeiros, 2023a). But, progress has been inconsistent. The report shows that respondents who are unemployed, have lower educational attainment, or earn less than 500 euros per month exhibit substantially lower resilience scores, even after controlling for other demographic characteristics through an OLS specification with robust standard errors. These descriptive findings establish the broad contours of vulnerability in the Portuguese population but leave unexplored whether the same variables, in combination with the richer set of behavioural and attitudinal items in the survey, can be used to build accurate individual-level predictions.

2.2 Financial Literacy and Its Relationship to Financial Vulnerability

2.2.1 The Three-Component Framework

Financial literacy is most commonly understood through the lens of the OECD/INFE framework, which decomposes it into three dimensions: knowledge, behaviours and attitudes (Atkinson & Messy, 2012). Financial knowledge refers to the ability to understand fundamental concepts such as interest rates, inflation, compound interest and risk diversification. Financial behaviours encompass actions like budgeting, saving regularly, keeping track of expenses and making informed product choices. Financial attitudes capture dispositions toward money, including preferences for spending versus saving and willingness to set long-term goals.

2.2.2 Knowledge, Behaviour and Outcomes

The relation between financial knowledge and economic outcomes has been widely studied, particularly since the foundational contributions of Lusardi and Mitchell (2014). Their analysis found strong associations between financial literacy, measured through the “Big Three” questions on compound interest, inflation and diversification, and a range of outcomes including retirement planning, stock-market participation and debt management. van Rooij et al. (2011) showed, using Dutch panel data, that individuals with lower financial knowledge were significantly less likely to invest in equities, even after controlling for risk aversion and income. The implication is that knowledge deficits may lead to suboptimal portfolio choices and, by extension, lower wealth accumulation over time.

However, the literature is far from unanimous on the causal direction of this association. Fernandes et al. (2014) conducted a meta-analysis of 168 papers and found that the effect of financial literacy interventions on downstream financial behaviours was, in most cases, rather small. They argued that interventions which target specific behaviours at the moment of decision tend to be more effective than general knowledge-building programmes. This finding is relevant for the present study because it raises the possibility that knowledge, as measured by quiz-type questions, may be less predictive of actual financial resilience than what people do with their money on a day-to-day basis.

2.2.3 The Portuguese Evidence

Research on financial literacy in Portugal has been constantly growing since the first wave of the BdP survey in 2010. Abreu and Mendes (2010) found that financially literate Portuguese investors held more diversified portfolios, consistent with the international evidence. More recently, the 2023 survey report itself contains detailed cross-tabulations showing that knowledge scores are positively correlated with income, education and employment status (Conselho Nacional de Supervisores Financeiros, 2023a). The average number of correct answers to the ten knowledge questions improved from 5.8 in 2020 to 6.4 in 2023, showing improvement over time, but still below the 2015 figure of 7.1.

The 2023 survey data shows that respondents who rated their own knowledge as “well above average” got seven or more items correct and only about a fifth of those who considered themselves “well below average” reached the same threshold. The association between self-assessment and objective performance is positive but far from perfect, and the residual variation may itself carry information about vulnerability (Conselho Nacional de Supervisores Financeiros, 2023a).

2.2.4 Financial Attitudes and the Behavioural Channel

While the bulk of the research literature has been devoted to information, attitudes and self-control are important determinants of financial results. Gathergood (2012) demonstrated that individuals with lower self-control were more prone to consumer over-indebtedness in the United Kingdom, even when accounting for differences in income and financial sophistication. The attitudinal items in the Portuguese survey, particularly the statements on spending orientation (B12 items 1 and 3), satisfaction with one’s financial situation (item 4) and goal-setting behaviour (item 7), tap into a similar set of psychological dispositions (Conselho Nacional de Supervisores Financeiros, 2023b).

It is possible, and the results of this dissertation will speak to this directly, that the behavioural and attitudinal dimensions of financial literacy matter more for prediction of vulnerability than objective knowledge. Someone who understands compound interest but fails to save, or sets no financial goals despite being capable of doing so, may be just as exposed to a shock as

someone with lower numeracy skills. The relative predictive contribution of these three dimensions is not, to the best of the author’s knowledge, well established in the Portuguese context, which is one of the contributions this study aims to make.

2.3 Machine Learning Applied to Financial Behaviour Prediction

2.3.1 Prior Applications in Household Finance

Machine learning techniques have been widely adopted in the financial industry, especially for credit scoring, fraud detection and default prediction (Leo et al., 2019). Khandani et al. (2010) demonstrated that combining traditional credit-bureau data with bank transaction records through tree-based ensemble methods substantially improved the prediction of consumer credit delinquency, relative to logistic regression baselines. Moscatelli et al. (2020) reached a similar conclusion, random forests and gradient boosted trees outperformed logistic regression in discriminatory power.

Applications specifically to survey-based measures of financial fragility or literacy are considerably less common. Most of the large-scale financial literacy studies cited above rely on standard econometric specifications, typically logit or probit models, sometimes augmented with instrumental variables to address endogeneity concerns. This disciplinary convention has left a methodological gap: the very surveys that collect rich batteries of behavioural and attitudinal items are rarely subjected to the kind of flexible, nonparametric modelling that could reveal which combinations of responses are most strongly associated with adverse outcomes.

2.3.2 Why Machine Learning May Add Value in This Context

The potential advantages of supervised classification methods over conventional regression for the task at hand can be stated concisely. Survey data of the kind used here are characterised by a large number of categorical and ordinal variables, many of which may interact in complex ways that are difficult to specify *a priori*. Logistic regression requires the analyst to decide in advance which interactions and polynomial terms to include; tree-based methods discover them automatically from the data (Hastie et al., 2009). A respondent who saved in the past year, holds a retirement product and reports being satisfied with their finances is plausibly different, in terms of vulnerability, from one who saved but holds no products and reports dissatisfaction. Logistic regression will capture each of those effects additively unless the analyst manually constructs interaction terms; a decision tree will split on whatever combination is most discriminative.

Nested cross-validation, in which an outer loop estimates generalisation performance while an inner loop tunes hyperparameters, is among the most rigorous approach for this purpose

(Cawley & Talbot, 2010; Varma & Simon, 2006). It avoids the optimistic bias that arises when the same data fold is used both to select a model and to evaluate it.

2.3.3 Key Algorithms

The six algorithms employed in this dissertation span a range of complexity and interpretability. Logistic regression serves as the parametric baseline and offers the advantage of directly interpretable coefficients. A single decision tree provides the most transparent non-linear alternative but is well known to suffer from high variance and sensitivity to small perturbations in the training data (James et al., 2021). Random forests address that instability by averaging the predictions of many decorrelated trees, each trained on a bootstrap sample of the data and a random subset of features (Breiman, 2001). K-nearest neighbours (KNN) is a non-parametric distance-based classifier that makes no distributional assumptions but can struggle in high-dimensional feature spaces without careful feature scaling and selection.

The two remaining models belong to the boosting family. Gradient boosting machines (GBM) build an ensemble of weak learners sequentially, with each new tree correcting the residual errors of its predecessors (Friedman, 2001). XGBoost extends this idea with regularisation terms that penalise tree complexity, along with engineering innovations such as column subsampling and efficient handling of missing values (Chen & Guestrin, 2016). In numerous applied competitions and benchmarking exercises, XGBoost and related gradient-boosted frameworks have emerged among the top-performing algorithms for tabular data. Whether that generalisation holds for a modestly sized survey dataset, where the risk of overfitting is higher than in the large-scale industrial datasets on which these algorithms were originally benchmarked, is one of the questions this dissertation addresses.

2.4 Research Gaps and Positioning of This Study

The preceding review suggests three gaps that this study is well positioned to address. The first is methodological: machine learning classifiers have not, to the author’s knowledge, been systematically applied to the Portuguese financial literacy survey data with the explicit goal of predicting financial vulnerability. Doing so not only tests whether flexible models outperform logistic regression in this specific setting but also yields a ranking of feature importance that can be compared with the findings of the descriptive analyses in the survey report.

The second gap is conceptual. Much of the existing literature treats financial knowledge, measured by quiz scores, as the primary driver of financial outcomes. The analysis presented here departs from that emphasis by including a comprehensive set of behavioural and attitudinal variables alongside knowledge items, and by letting the models determine empirically which dimension carries the most predictive weight. If behaviours and attitudes dominate, the implication for financial education policy is significant: interventions that focus exclusively on

teaching people what compound interest is may be less effective than programmes that help them change what they do with their money.

The third gap concerns the characterisation of misclassified individuals. Most classification studies report aggregate performance metrics and, at best, confusion matrices. This dissertation goes a step further by profiling the vulnerable respondents the best model classified as resilient. Understanding who falls through the cracks of a predictive model is valuable both for refining the model and for designing safety nets that reach the right people.

The chapter that follows describes the data, the feature engineering pipeline and the modelling framework in detail, laying the methodological groundwork for the empirical analysis.

3. Data and Methodology

This chapter describes the dataset, the construction and operationalisation of the target variable, the feature engineering and preprocessing steps, and the model training and evaluation framework. The goal is to document each methodological decision with enough precision that the analysis could be replicated by another researcher with access to the same microdata.

3.1 Data Source

The data used in this dissertation comes from the 4th Survey on the Financial Literacy of the Portuguese Population, done in collaboration by Banco de Portugal, CMVM and ASF under the auspices of the National Plan for Financial Education. The data collection was carried out by Eurosondagem between 6 January and 13 February 2023 (Conselho Nacional de Supervisores Financeiros, 2023a). A total of 1,510 face-to-face interviews were conducted across all seven NUTS II regions of Portugal, targeting the resident population aged 16 and above. The sample was drawn using a random-route method and stratified by gender, age, geographic location, employment situation and schooling level, with quotas calibrated to the proportions observed in the 2021 Census published by the Instituto Nacional de Estatística. The sample carries an average margin of error of 2.5 percentage points at the 95 per cent confidence level. The microdata file contains 1,513 records, three more than the 1,510 completed interviews cited in the report. Seven were identified as unsuitable for analysis: four were empty across all fields, and three reported an age of 999, clearly not a plausible value. Removing these left a working sample of 1,506 observations.

The questionnaire follows the OECD/INFE Toolkit for Measuring Financial Literacy and Financial Inclusion (OECD, 2022) and is organised into four sections. Section A collects demographic and socio-economic characteristics including region, gender, age, education, employment status, nationality and household composition. Section B covers personal finance planning, savings behaviour, retirement confidence and attitudinal statements. Section C records awareness and holdings of financial products. Section D measures financial knowledge through numeracy questions, true/false conceptual items, self-assessed knowledge, digital literacy and household income.

3.2 Target Variable

The binary outcomes used throughout the analysis are derived from question B6, which asks: “If you personally faced a major expense today, equivalent to your own monthly income, would you be able to pay it without borrowing the money or asking family or friends to help?” Re-

spondents who answered “Yes” are coded as *Resilient* (1); those who answered “No” are coded as *Vulnerable* (0).

Three response categories were excluded prior to modelling, accounting for 133 respondents, or 8.8 per cent of the 1,506 usable records. The largest group consists of 109 respondents who selected “Not applicable”, and these were removed on definitional grounds. Because question B6 frames the shock as an expense equivalent to the respondent’s own monthly income, the item is not applicable for individuals without personal income. The remaining 24 cases are genuine non-responses, comprising 19 “Do not know” and 5 “Refused to answer” answers. Neither can be assigned to the resilient or the vulnerable class without imputation assumptions that would be difficult to defend for a binary outcome derived from a single item.

One implication of this filter should be stated directly. Because the “Not applicable” exclusion removes individuals who report no personal income, the analytical population is effectively restricted to respondents with an income of their own, and the findings are best read as applying to that population. This qualification is mentioned again in Section 5.3.

Two related variables were excluded from the feature set to prevent data leakage. B7 asks whether the respondent’s income was insufficient to cover living costs in the past year, and B9 asks how long the respondent could continue to cover living expenses if they lost their main income source. Both are near-direct proxies for the target and would artificially inflate model performance if retained.

3.3 Feature Engineering

The raw questionnaire contains 229 variables spanning the four survey sections. Several engineering steps were applied to transform these into a format suitable for supervised classification.

3.3.1 Knowledge Score Construction

A composite `knowledge_score` variable was constructed by summing correct responses to questions D2 through D7. Questions D2–D6 each return one binary indicator covering, respectively, basic division, inflation understanding, interest on a one-day loan, simple-interest calculation and compound interest. Question D7 contributes six binary sub-items: three from the standard OECD/INFE financial knowledge module (high inflation raises the cost of living; high returns imply high risk; diversification reduces investment risk) and three from the digital financial literacy extension (whether digital contracts require a paper signature; whether sharing personal data enables targeted offers; whether cryptocurrencies carry the same legal-tender status as cash). Each correctly answered sub-item contributes one point.

3.3.2 Multicoded Variables

Several questions allow multiple responses. B4, for instance, asks about budgeting and financial management practices and permits the respondent to indicate all that apply. B5 records savings methods, and B11 records expected sources of retirement funding. Each response option in these multicoded questions was converted into a separate binary indicator (1 if selected, 0 otherwise), following standard survey-data practice. A variable indicating whether the respondent refused to answer was also created.

3.3.3 Ordinal Encoding

Inherently ordered variables such as household income range (D11), retirement planning confidence (B10) and self-assessed knowledge (D1) were encoded numerically while preserving their ordinal structure. For D11, the income brackets were mapped to integers from 1 (up to 500 euros) through 4 (above 2,500 euros), with refusals treated as missing. Education level (A4) was similarly mapped onto an integer scale from 0 (no formal education) to 5 (postgraduate), following the ordering in the questionnaire.

3.4 Data Cleaning

After applying the B6 exclusion filter and engineering the variables described above, the dataset contained 1,373 observations and 149 columns. Rows with any remaining missing values were dropped, yielding a final analytical sample of 1,355 observations. This listwise deletion approach was chosen over imputation because the missing values in this dataset are not randomly distributed but tend to cluster among respondents who refused to answer income or product-holding questions, a pattern that complicates the assumptions required by standard imputation methods.

The three records reporting an age of 999, noted in Section 3.1, were treated as data-entry errors and removed during the initial inspection. Beyond this quality filter, no statistical outlier treatment was applied to the age variable (A3), although a small number of respondents aged 80 and above were retained in the analysis despite being at the tail of the distribution. The decision not to trim age reflects the fact that elderly respondents constitute a substantively important group for financial vulnerability research, and excluding them would compromise the representativeness of the sample.

3.5 Modelling Pipeline

The data were split into training and test sets using a stratified 80/20 random partition, preserving the class distribution in both subsets. The training set contained 1,084 observations (747 resilient, 337 vulnerable) and the test set contained 271 (187 resilient, 84 vulnerable).

3.5.1 Preprocessing

Separate preprocessing pipelines were applied to numeric and categorical features. The numeric pipeline consisted of four sequential steps. First, a custom `DropCorrelatedColumns` transformer (Pearson threshold $r > 0.8$) removed one member of each highly correlated feature pair; this step is applied before scaling so that correlations are computed on the original numeric scale. Second, a median imputer filled any residual missing values. Third, a `VarianceThreshold` filter (threshold = 0.08) discarded near-constant features that carry negligible predictive signal. Fourth, `StandardScaler` centred and scaled each remaining feature to zero mean and unit variance, as required by the distance-based and regularised models (KNN and logistic regression). The rationale for correlation filtering is that retaining highly collinear features can distort coefficient estimates in logistic regression and inflate importance scores in tree-based models without providing additional discriminative information.

The categorical pipeline applied a `RelativeFreqGrouper` that collapsed categories appearing in fewer than 5 per cent of the training data into a single “infrequent” label, followed by one-hot encoding. This grouping step prevents the creation of extremely sparse dummy columns that would contribute noise rather than signal, and it mitigates the risk of overfitting to rare categories that may not appear in the test set.

After preprocessing, the training feature matrix contained 126 columns: 106 numeric and 20 categorical (post-encoding).

3.5.2 Class Imbalance Handling

The training sample is moderately imbalanced, with roughly 69 per cent of cases resilient and 31 per cent vulnerable. A common response is to oversample the minority class using SMOTE (Chawla et al., 2002). That route was avoided here. SMOTE builds new minority cases by interpolating between existing ones, which amounts to inventing survey respondents who never sat the questionnaire, and that is difficult to defend when every feature is a recorded answer. Imbalance was handled inside the models instead. Logistic regression, the decision tree and the random forest were fitted with `class_weight='balanced'`, so that errors on the vulnerable class weigh more during training. XGBoost has no equivalent argument, so its `scale_pos_weight` was set to the ratio of negative to positive cases ($337/747 \approx 0.45$) to the same end (Chen & Guestrin, 2016). The gradient boosting and KNN implementations offer no comparable setting and were left unadjusted;

3.6 Model Selection and Evaluation

3.6.1 Nested Cross-Validation

Model selection and hyperparameter tuning were conducted using a nested cross-validation design. The outer loop used a 5-fold stratified split to estimate generalisation performance. Within each outer fold, an inner loop performed randomised hyperparameter search (`RandomizedSearchCV`) over 30 candidate configurations, also using 5-fold stratified cross-validation, with ROC-AUC as the scoring criterion. This nested structure ensures that the performance estimates reported in the outer loop are not contaminated by the hyperparameter tuning process, a concern that is well documented in the methodological literature (Cawley & Talbot, 2010; Varma & Simon, 2006).

Six algorithms were evaluated: logistic regression, decision tree, random forest, K-nearest neighbours, gradient boosting and XGBoost. For each algorithm, the hyperparameter search space was defined to cover a range of regularisation strengths, tree depths, ensemble sizes and learning rates appropriate to the model family. The specific grids are reported in Appendix C.

3.6.2 Evaluation Metrics

Given the moderate class imbalance and the policy relevance of correctly identifying vulnerable individuals, several complementary metrics were used. ROC-AUC serves as the primary comparison criterion because it evaluates discriminative ability across all classification thresholds. Balanced accuracy, defined as the arithmetic mean of sensitivity (recall) and specificity, accounts for unequal class sizes. Precision, recall, F1-score and the full confusion matrix are also reported for the test set.

For the machine-learning metrics in this dissertation, the *positive* class is the resilient one (B6 = Yes). This choice follows scikit-learn’s convention of treating the label coded 1 as positive, and it is the basis on which precision, recall and the F1-score are reported throughout. PR-AUC (the area under the precision-recall curve) is informative chiefly when the positive class is the minority (Saito & Rehmsmeier, 2015). Here the positive class is the majority (approximately 69 per cent of observations), so the PR-AUC values reported in the next chapter should be read as discrimination of the resilient majority rather than of the vulnerable minority; even a trivial classifier that always predicts “resilient” would achieve a non-trivial baseline PR-AUC equal to the prevalence of that class. The figure is included for completeness and for comparability with the existing literature, but ROC-AUC and balanced accuracy carry the bulk of the interpretive weight. The same reasoning applies to the F1-score reported for the resilient class: because it combines recall with a precision whose floor equals the positive-class prevalence, a trivial all-resilient classifier would already reach an F1 of roughly 0.82, so the F1 values reported in the next chapter are read as corroborating rather than decisive.

The decision to use ROC-AUC rather than accuracy as the primary metric reflects the fact that a naive classifier predicting “resilient” for every observation would achieve roughly 69 per cent accuracy, a misleadingly high number that masks total failure to identify the vulnerable class.

4. Results

This chapter presents the empirical findings in four parts. It begins with a summary of the exploratory data analysis, moves to the model performance comparison, examines the best model in detail and concludes with an assessment of the key predictors of financial vulnerability.

4.1 Exploratory Data Analysis

4.1.1 Target Distribution

After applying the B6 exclusion filter and dropping rows with missing values, the final sample contained 1,355 observations. Of these, 934 (68.9 per cent) reported being able to cover an unexpected expense equivalent to one month of income without borrowing, and 421 (31.1 per cent) reported that they could not. The 80/20 stratified train-test split preserved these proportions closely: the training set contained 68.9 per cent resilient and 31.1 per cent vulnerable cases, and the test set 69.0 per cent and 31.0 per cent, respectively. Figure 4.1 displays the class distribution.

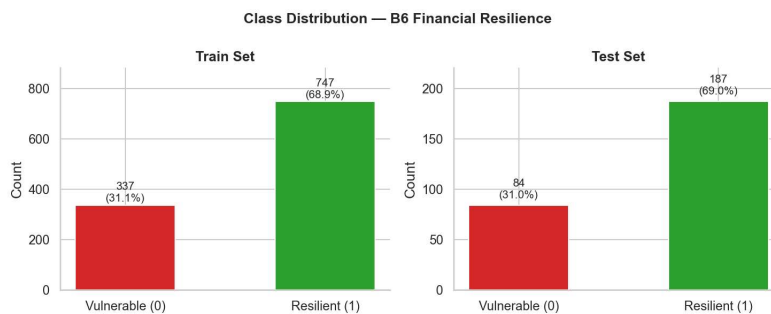


Figure 4.1. Distribution of the target variable (B6) in the final analytical sample after exclusions and listwise deletion of missing values. The resilient class accounts for approximately 69 per cent of observations.

4.1.2 Demographic and Economic Patterns

Several demographic and economic variables demonstrate a clear univariate associations with financial vulnerability, consistent with the patterns documented in the survey report (Conselho Nacional de Supervisores Financeiros, 2023a). Respondents in the lowest income bracket (up to 500 euros per month) are substantially more likely to report inability to absorb a shock than those earning above 2,500 euros. Education follows a similar gradient: vulnerability rates decline as educational attainment increases, with the sharpest drop occurring between basic education (9th year or below) and secondary education. Unemployed respondents and those clas-

sified as “non-active, other situation” show the highest vulnerability rates among employment categories, while retirees occupy an intermediate position.

Age shows a clear gradient, and it runs in the direction one might naively expect rather than against it. The youngest respondents (16 to 24) are by some margin the most vulnerable group, with over half reporting that they could not cover a one-month income shock. Vulnerability then falls across the working years. The 25 to 39 and 40 to 54 brackets are the most resilient, before the rate climbs again, more gently, among the 55 to 69 and 70-plus groups. What makes the young-adult figure worth focusing on is that it survives the sample construction rather than being produced by it. Because respondents without personal income were removed by the “not applicable” filter, the 16-to-24-year-olds who remain are those already earning; even this comparatively advantaged slice records the highest vulnerability rate in the sample. The reading consistent with this is one in which young earners combine low and often irregular incomes with little accumulated saving, which leaves a thin buffer against an unexpected cost.



Figure 4.2. Financial resilience (ability to cover an unexpected expense equivalent to one month’s personal income) by gender, age band, education level and region of residence. Bars show the within-group share of resilient (Yes) and vulnerable (No) respondents in the analytical sample (n = 1,355). Percentages are row-wise within each category.

4.1.3 Knowledge Score Distribution

The composite knowledge score, constructed from questions D2 through D7 as described in Section 3.3.1, has a roughly bell-shaped distribution with a mean of approximately 6.6 points out of a maximum of eleven. Resilient respondents score moderately higher on average than vulnerable ones, but the distributions overlap considerably, suggesting that knowledge alone is unlikely to be a strong discriminator. Figure 4.3 shows the distribution of the knowledge score, broken down by target class.

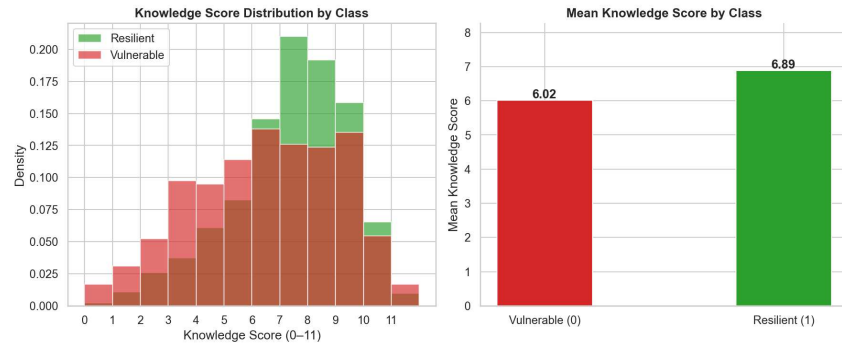


Figure 4.3. Distribution of the composite financial knowledge score (D2–D7) by target class. While resilient respondents score somewhat higher on average, the overlap between the two distributions is substantial.

4.2 Model Performance Comparison

4.2.1 Nested Cross-Validation Results

Table 4.1 reports the mean and standard deviation of ROC-AUC across the five outer folds for each model. XGBoost achieved the highest mean (0.896), followed by random forest (0.887) and gradient boosting (0.883). Logistic regression and KNN were close behind, at roughly 0.867 and 0.866. The single decision tree was the weakest (0.825). Standard deviations ranged from 0.023 to 0.030, so no model was dramatically unstable across folds.

Figure 4.4 shows these distributions as box plots.

Table 4.1. Nested cross-validation results (outer 5-fold, inner RandomizedSearchCV with 30 iterations and 5-fold CV, scored on ROC-AUC). Models are sorted by mean ROC-AUC.

Model	Mean ROC-AUC	Std. Dev.
XGBoost	0.896	0.023
Random Forest	0.887	0.025
Gradient Boosting	0.883	0.025
Logistic Regression	0.867	0.029
KNN	0.866	0.030
Decision Tree	0.825	0.024

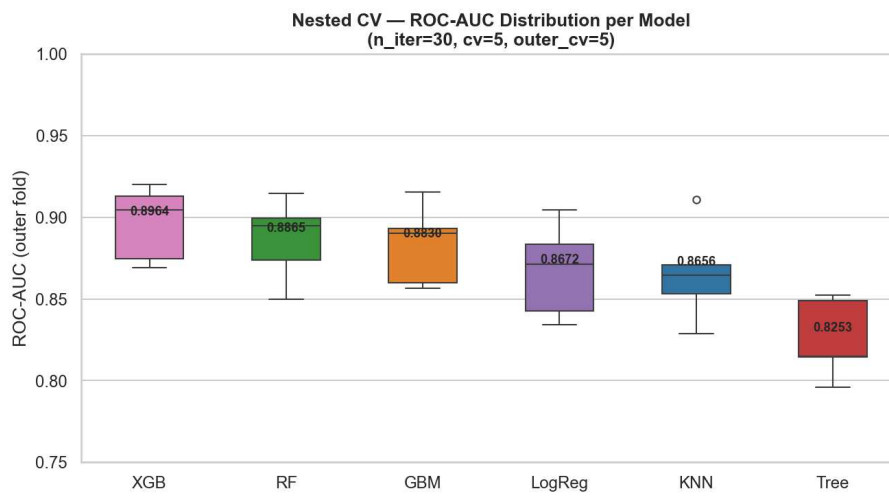


Figure 4.4. Box plot of ROC-AUC scores across the five outer cross-validation folds for each model. XGBoost shows both the highest median and the tightest interquartile range.

4.2.2 Test Set Performance

Table 4.2 reports the full set of evaluation metrics on the held-out test set ($n = 271$). XGBoost again ranks first, with a ROC-AUC of 0.913 and a PR-AUC of 0.954, confirming the ordering obtained from the nested cross-validation. Its balanced accuracy of 0.833 reflects an even trade-off between recall (0.893) and specificity (0.774). Gradient boosting and random forest follow closely, with test ROC-AUC values of 0.909 and 0.906.

The decision tree and logistic regression form a contrasting group. The tree records the highest precision (0.918) and specificity (0.845) of any model, but also the lowest recall (0.781). Logistic regression behaves similarly, ranking second on both precision (0.905) and specificity (0.810). Both classify conservatively: they label a respondent as resilient only when the evidence is strong, which keeps their false positives low but causes them to miss more of the resilient cases. This is the familiar precision-recall trade-off.

KNN sits at the opposite extreme. It has the highest recall (0.925) but the lowest specificity (0.667), meaning it labels most respondents as resilient regardless of their true status.

Table 4.2. Model performance on the held-out test set (n = 271). The positive class is Resilient (1). Models are sorted by ROC-AUC.

Model	Acc.	Prec.	Rec.	Spec.	Bal. Acc.	F1	ROC-AUC	PR-AUC
XGB	0.856	0.898	0.893	0.774	0.833	0.895	0.913	0.954
GBM	0.841	0.875	0.898	0.714	0.806	0.887	0.909	0.949
RF	0.841	0.891	0.877	0.762	0.820	0.884	0.906	0.949
KNN	0.845	0.861	0.925	0.667	0.796	0.892	0.901	0.947
LogReg	0.812	0.905	0.813	0.810	0.811	0.856	0.897	0.952
Tree	0.801	0.918	0.781	0.845	0.813	0.844	0.866	0.902



Figure 4.5. Test-set ROC-AUC for all six models, presented as a bar chart. All ensemble methods exceed 0.900, with XGBoost narrowly leading. The chart complements the cross-validation distribution shown in Figure 4.4.

4.2.3 ROC and Precision-Recall Curves

Figure 4.6 presents the ROC curves for the six models on the test set. The curves cluster tightly in the upper-left region, confirming that all models discriminate well. XGBoost's curve lies slightly above the rest at most operating points, consistent with its higher ROC-AUC. Figure 4.7 shows the precision-recall curves. Here logistic regression's PR-AUC (0.952) almost matches XGBoost's (0.954), so the two are close in their ability to rank resilient respondents correctly, even if XGBoost remains the more balanced model across the full range of thresholds.

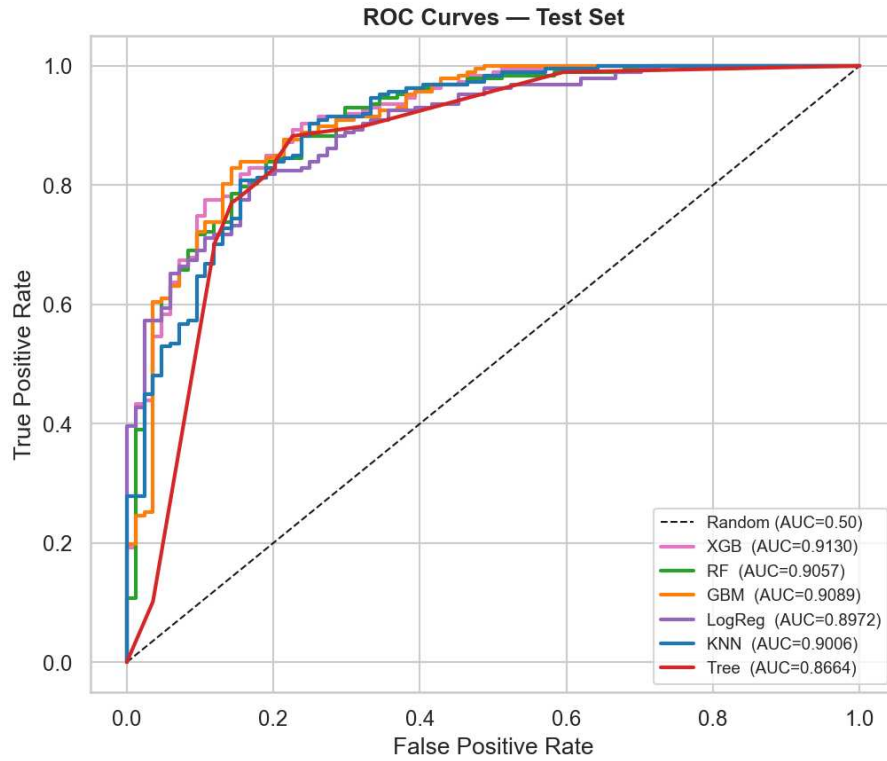


Figure 4.6. Receiver operating characteristic (ROC) curves for all six models on the held-out test set. The diagonal represents random guessing.

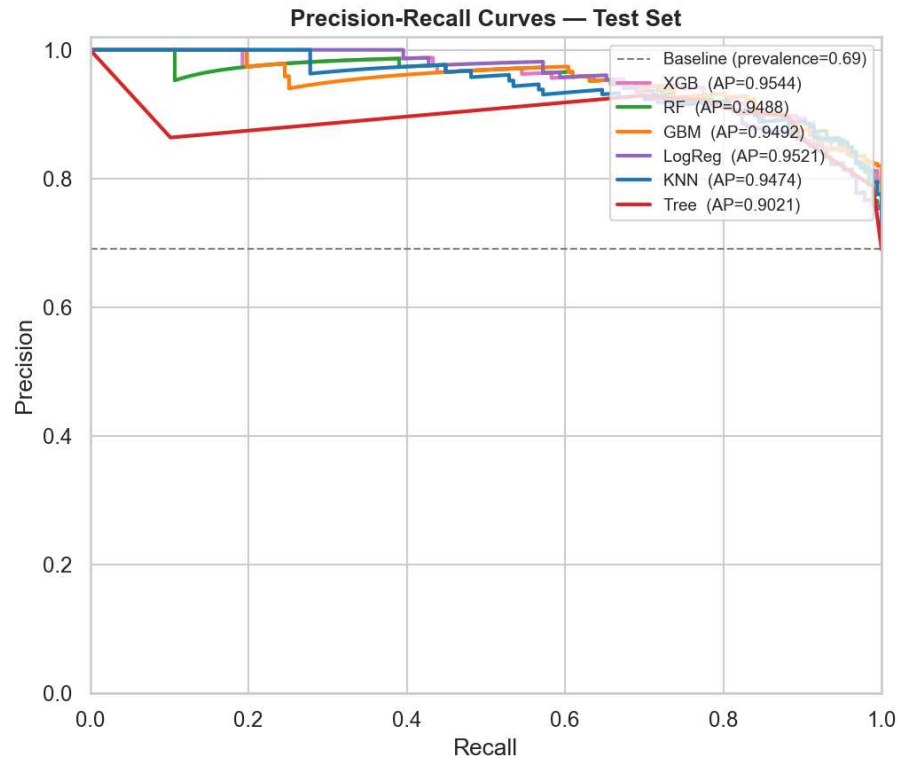


Figure 4.7. Precision-recall curves for all six models on the test set. Logistic regression and XGBoost achieve near-identical PR-AUC values despite differing in specificity and recall trade-offs.

4.2.4 Overfitting Assessment

Table 4.3 reports the gap between training and test ROC-AUC for each model. A large gap signals that the model has memorised training-set noise rather than learning generalisable patterns. KNN shows the largest gap (0.099), but its training ROC-AUC of 1.000 is an artefact of the weighting scheme, not a sign of strong learning: when the inner cross-validation selects `weights='distance'` (see Appendix C), each training point predicts its own label, so the training score collapses to a perfect value. Its test AUC of 0.901 is the more meaningful figure. The decision tree has the smallest gap (0.008), reflecting its constrained complexity after pruning. Random forest (0.042) and logistic regression (0.037) generalise most cleanly, while XGBoost (0.073) and gradient boosting (0.072) show moderate gaps that remain acceptable for a dataset of this size. Figure 4.8 visualises these gaps.

Table 4.3. Train versus test ROC-AUC gap as an indicator of overfitting. Smaller gaps indicate better generalisation.

Model	Train ROC-AUC	Test ROC-AUC	Gap
XGB	0.986	0.913	0.073
GBM	0.981	0.909	0.072
RF	0.948	0.906	0.042
KNN	1.000	0.901	0.099
LogReg	0.934	0.897	0.037
Tree	0.875	0.866	0.008

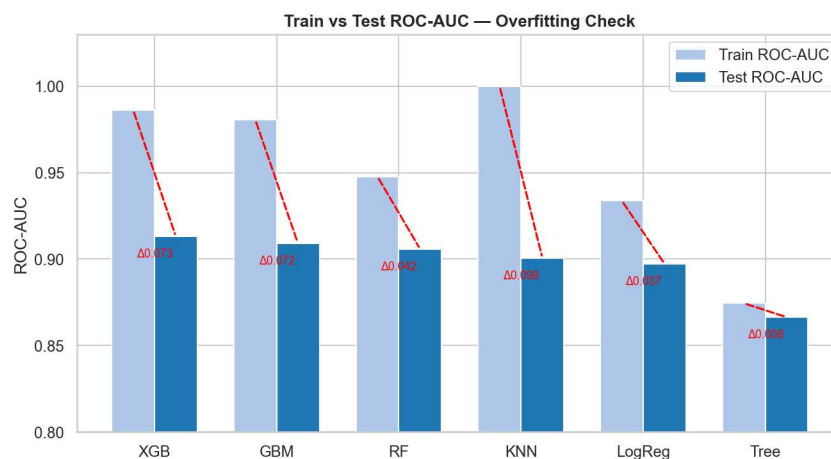


Figure 4.8. Training versus test ROC-AUC gap for each model. The KNN gap reflects an artefact of distance-based weighting on the training set rather than genuine in-sample learning; among ensemble models, random forest generalises most cleanly.

4.3 Feature Importance and Model Interpretability

XGBoost emerged from the previous section as the strongest model, but its built-in importance scores are awkward to interpret: the same model yields different rankings depending on whether importance is measured by gain, by split frequency or by cover, and the values can also shift when correlated columns are dropped during preprocessing. The discussion of *which* features drive predictions is therefore organised around two more transparent vantage points: impurity-based importances from the random forest (Subsection 4.3.1) and signed coefficients from the regularised logistic regression (Subsection 4.3.2). Both point to the same set of dominant features, but they make the pattern easier to read. Confusion matrices for all six models then close the section (Subsection 4.3.3) and motivate the error analysis that follows.

4.3.1 Feature Importance: Random Forest

Random forest importances are reported here, rather than XGBoost's, because they rest on a single well-understood impurity-reduction criterion and so read more cleanly as a ranking. Table 4.4 lists the sixteen most important features. The strongest single predictor is B5_Did_not_save_last_year which alone accounts for just over 10 per cent of total importance. Retirement planning confidence (B10) ranks second, followed by household monthly income range (D11) and whether the respondent left savings in a current account (B5). Figure 4.9 shows the full top-20 ranking.

Table 4.4. Random forest feature importance (top 16 by impurity-based importance). Behavioural and attitudinal variables dominate the ranking; the composite knowledge score appears at position 16.

Rank	Feature	Importance
1	B5: Did not save last year	0.102
2	B10: Retirement planning confidence	0.072
3	D11: Household monthly income range	0.060
4	B5: Left money in current account	0.057
5	B12: Financial situation limits goals (Agree)	0.047
6	B12: Satisfied with finances (Agree)	0.047
7	B11: Savings used for retirement	0.039
8	C1.2: Term deposits (Holds)	0.032
9	B12: Financial situation limits goals (Disagree)	0.032
10	B4: Has automatic payments	0.028
11	D1: Self-assessed financial knowledge	0.026
12	A3: Age	0.024
13	C1.2: Retirement savings plans (Holds)	0.019
14	C1.2: Insurance products (Holds)	0.017
15	B12: Satisfied with finances (Disagree)	0.015
16	Knowledge score (D2–D7 composite)	0.014

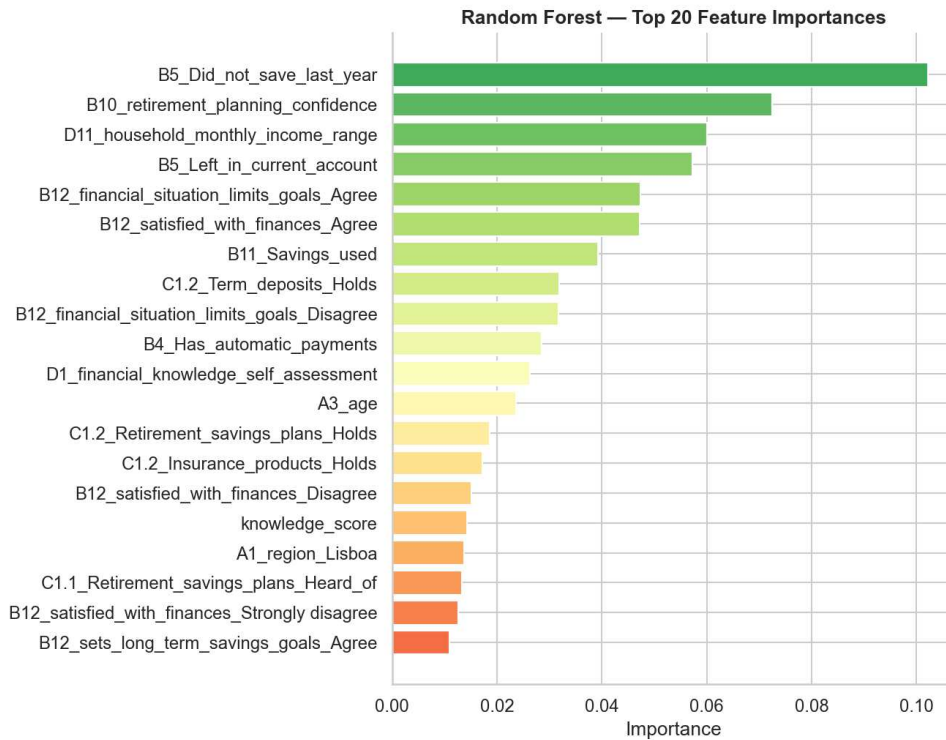


Figure 4.9. Top 20 features by random forest impurity-based importance. Savings behaviour (B5), retirement confidence (B10) and income (D11) clearly outweigh objective knowledge scores.

A striking finding is the low ranking of the composite knowledge score. At position 16 with an importance of 0.014, it contributes roughly one-seventh of what the leading behavioural variable contributes. Self-assessed knowledge (D1) ranks higher, at position 11, which is itself noteworthy: respondents' own perception of their financial understanding appears to carry more predictive signal than their actual quiz performance. This result is consistent with the hypothesis, developed in the literature review, that what people *do* with their money matters more for vulnerability than what they *know* about financial concepts.

4.3.2 Logistic Regression Coefficients

The logistic regression coefficients offer an interpretable complement to the random forest ranking. Figure 4.10 shows the fifteen largest coefficients by absolute magnitude. Positive coefficients are associated with a higher predicted probability of resilience and negative coefficients with a lower one. Several B12 attitudinal items appear among the largest coefficients, alongside savings and goal-setting indicators, broadly echoing the variables that rank highly in the random forest. The two rankings are produced by different methods and need not agree exactly, but the overlap in the variables they emphasise lends some confidence to the overall picture.

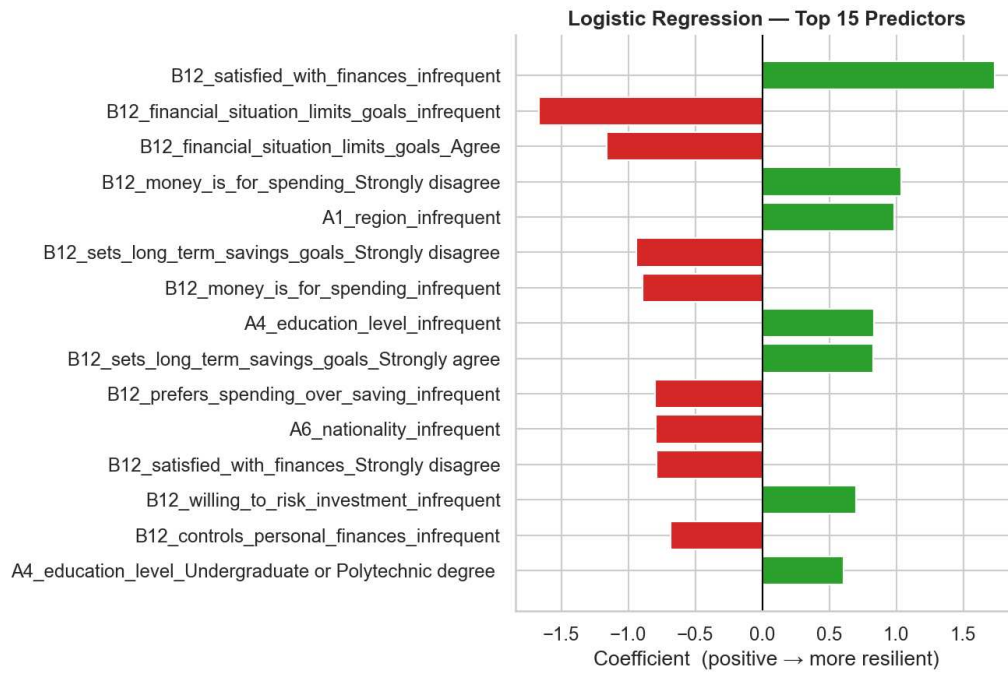


Figure 4.10. Top 15 logistic regression coefficients sorted by absolute magnitude. Positive coefficients indicate a higher probability of being classified as resilient; negative coefficients indicate higher vulnerability.

4.3.3 Confusion Matrices

Figure 4.11 presents the confusion matrices for all six models on the test set. Following the ML-standard convention adopted throughout this dissertation, the positive class is the resilient one (B6 = Yes). On that convention, XGBoost correctly identifies 167 of 187 resilient cases and 65 of 84 vulnerable cases, yielding 20 false negatives (resilient respondents classified as vulnerable) and 19 false positives (vulnerable respondents classified as resilient). The latter group, the 19 vulnerable individuals misclassified as resilient, is the focus of the profile analysis in Section 4.4. By comparison, logistic regression achieves fewer false positives (16) but substantially more false negatives (35), confirming its conservative tendency with respect to the resilient class. The decision tree produces the fewest false positives (13) but the most false negatives (41), reflecting the tightest classification boundary among the six models.

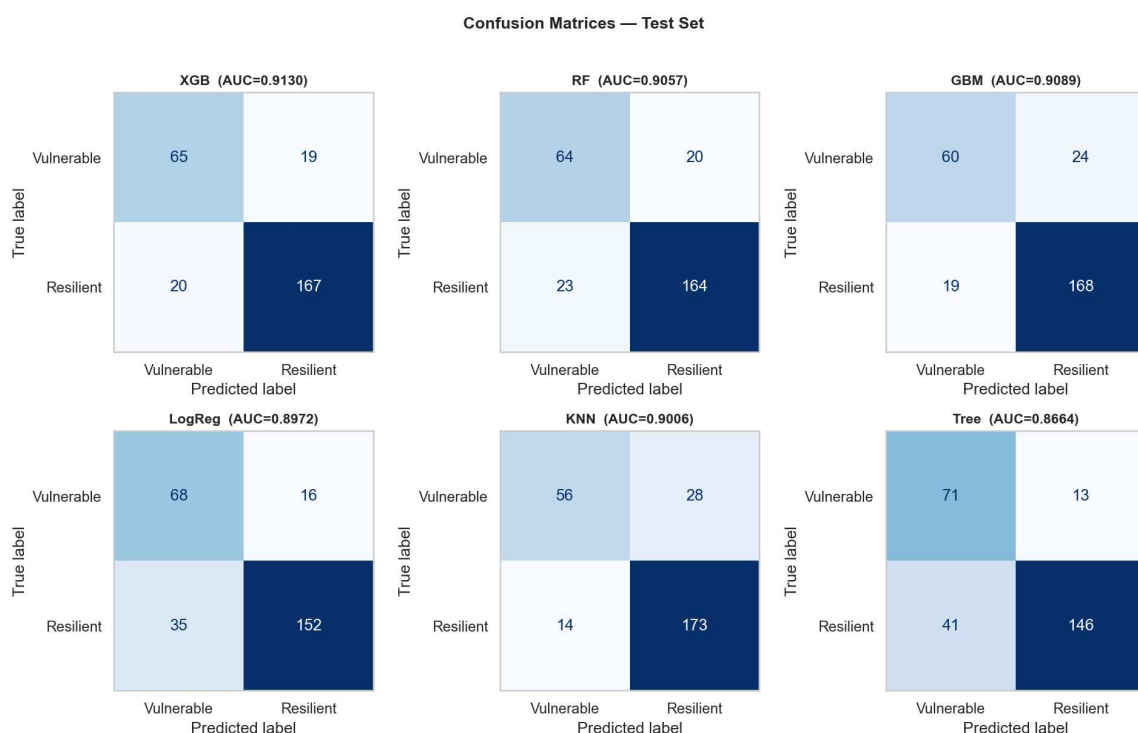


Figure 4.11. Confusion matrices for all six models on the held-out test set ($n = 271$). Each cell reports the absolute count; row labels indicate the true class and column labels the predicted class.

4.4 Profile of Vulnerable Respondents Classified as Resilient

From a policy standpoint, the most consequential errors are not the cases where the model wrongly labels a resilient respondent as vulnerable, which would mainly cause unneeded outreach, but the cases where a vulnerable respondent is wrongly labelled as resilient and is therefore missed by any targeting rule built on the model's output. Under the resilient-as-positive convention used in the confusion matrix, these are the 19 false positives reported for XGBoost in Section 4.3.3.

The 19 vulnerable respondents that XGBoost classified as resilient represent 22.6 per cent of the 84 truly vulnerable cases. Profiling these misclassified individuals reveals a coherent pattern. Table 4.5 compares the mean or modal values of key features across three groups: the 19 misclassified vulnerable respondents, the 65 vulnerable respondents the model correctly identifies (true negatives, TN) and the full test set.

Table 4.5. Profile of the 19 vulnerable respondents XGBoost classified as resilient, compared to the 65 correctly classified vulnerable respondents (TN) and the full test set. Numeric features report means; categorical features report modes.

Feature	Misclassified (n=19)	TN (n=65)	All Test
B5: Did not save last year	0.316	0.862	0.410
B5: Left in current account	0.632	0.108	0.487
B11: Savings used	0.421	0.015	0.303
B10: Retirement planning confidence	1.263	0.292	1.188
D11: Income range	2.211	1.615	2.520
C1.2: Retirement savings plans (Holds)	0.316	0.000	0.225
B12: Satisfied with finances	Agree	Neither	Agree
B12: Financial situation limits goals	Agree	Agree	Agree

These respondents look, on the surface, like resilient individuals. Only 31.6 per cent of them did not save in the past year, compared with 86.2 per cent of the correctly classified vulnerable cases. Nearly two-thirds left money in a current account, and 31.6 per cent hold retirement savings plans. Their average income is in the mid-range (bracket 2.2, corresponding roughly to 500–1,000 euros), and they report being satisfied with their financial situation. Yet they are unable to cover a one-month income shock.

One possible interpretation is that these individuals project an outward appearance of financial soundness. They are considerably more likely than the correctly classified vulnerable group to hold formal financial products, and they report moderate confidence in their retirement planning (1.26 on a 0–5 scale, against 0.29 for the TN group). The variable B11_Savings_used, which records whether a respondent expects to fund retirement by drawing on personal savings, is far more common among them (mean 0.42 for the misclassified group versus 0.015 for TN). Taken together, these signals describe people with a savings-oriented, largely self-reliant view of their own retirement who are nonetheless unable to cover a one-month income shock today. Their planning confidence, though higher than that of the clearly vulnerable group, remains well short of the levels associated with genuine financial security, and it is plausibly this combination of reassuring outward signals and hidden fragility that misleads the model.

Figure 4.12 visualises this profile across the top numeric features.

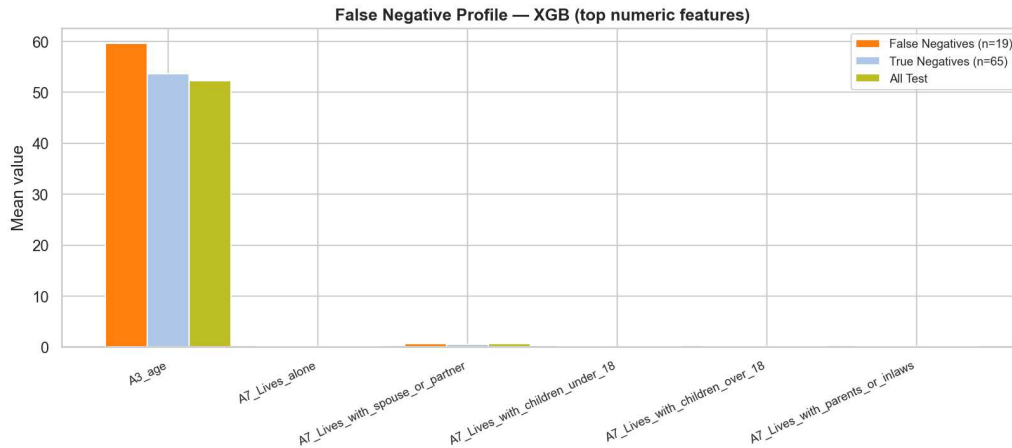


Figure 4.12. Comparison of mean feature values for the 19 vulnerable respondents misclassified by XGBoost as resilient, the 65 correctly classified vulnerable respondents (TN) and the full test set, across selected numeric features.

4.5 Key Predictors of Financial Vulnerability

Synthesising the evidence from the random forest importance ranking, the logistic regression coefficients and the error analysis, three broad categories of predictors emerge as the most influential.

The first and most powerful category is *savings behaviour*. Whether a respondent saved in the past year, and in what form, dominates both the RF importance ranking and the error profile. Not saving at all is the single strongest signal of vulnerability, while leaving money in a current account or a term deposit is associated with resilience. A related signal, the B11_Savings_used variable, records whether a respondent expects to fund retirement by drawing on personal savings. Its prominence in the error analysis is telling: a stated intention to rely on one's own savings in retirement does not guarantee the short-term liquidity needed to absorb an immediate shock, which helps explain why this group is so readily mistaken for resilient.

The second category is *financial attitudes and subjective assessments*. The B12 Likert items on financial satisfaction, perceived constraint and goal-setting behaviour appear consistently among the top predictors in both model families. Self-assessed financial knowledge (D1) also contributes meaningfully, ranking above the objective knowledge score in the RF importance hierarchy. These findings are consistent with a reading of financial vulnerability as a partly psychological phenomenon: individuals who feel constrained by their financial situation, who do not set goals or who express dissatisfaction with their finances are more likely to be vulnerable, independently of their income or knowledge level.

The third category is *household income and product holdings*. Income range (D11) is the third-ranked RF feature, and holding term deposits or retirement savings plans is associated with resilience. Age contributes at a moderate level (rank 12 in RF), consistent with the non-linear relationship observed in the EDA.

The composite knowledge score (D2–D7), by contrast, ranks 16th out of 126 features in the random forest. While it is not irrelevant, it contributes considerably less predictive power than savings behaviour, attitudes or income. This finding has direct implications for the design of financial education programmes, an issue taken up in the next chapter.

5. Discussion

The results reported in the previous chapter raise several questions that warrant careful interpretation. This chapter organises the discussion around four themes: the relationship between the empirical findings and the existing literature, the implications for financial education policy in Portugal, the limitations of the study, and directions for future research.

5.1 Interpretation of Results in Light of the Literature

The headline finding of this dissertation is that financial vulnerability, as captured by the inability to cover an unexpected one-month income shock, can be predicted with a test-set ROC-AUC of 0.913 using an XGBoost classifier trained on survey responses. This level of discrimination is high by the standards of social-science prediction tasks, where noise levels tend to be substantial. It falls comfortably within the range reported by Khandani et al. (2010) for consumer credit-risk models, although direct comparison is complicated by differences in sample size, feature granularity and target definition. What the result does suggest, unambiguously, is that the survey data contain sufficient predictive structure to support individual-level classification, not merely population-level description.

More consequential than the aggregate performance figure is the ranking of predictors. The literature on financial literacy has, for good reasons, placed considerable emphasis on the role of objective knowledge. The foundational contributions of Lusardi and Mitchell (2014) and van Rooij et al. (2011) established strong associations between quiz-based knowledge measures and financial outcomes ranging from stock-market participation to retirement planning. In this dissertation, however, the composite knowledge score constructed from questions D2 through D7 ranks sixteenth out of 126 features in the random forest importance hierarchy, contributing barely one-seventh of the predictive weight carried by the top-ranked variable, which is a simple binary indicator of whether the respondent saved in the past year.

This finding does not invalidate the importance of financial knowledge, it does, however, complicate the narrative in a way that is consistent with the more cautious strand of the literature. Fernandes et al. (2014) found, in their meta-analysis, that the effect of financial literacy on behaviour was often small once confounders were controlled for, and that interventions targeting specific behaviours outperformed generic knowledge-building programmes. The results presented here offer a complementary piece of evidence from a different methodological vantage point: when a flexible classifier is given access to knowledge, behavioural and attitudinal variables simultaneously, it relies primarily on the latter two.

One interpretation for self-assessed financial knowledge (D1) ranking higher than the objective score is that self-assessment captures actual knowledge and financial confidence that is more behaviourally relevant than quiz performance alone. A respondent who rates them-

selves as financially knowledgeable may, on average, be more likely to act on that knowledge, for instance by saving regularly or shopping for financial products. The concept of financial self-efficacy, though not explicitly measured in the survey, is a plausible mediating mechanism. Alternatively, the finding could partly reflect reverse causation: people who have managed to build a savings buffer may retrospectively rate their own knowledge more highly, precisely because their financial situation has turned out well. The cross-sectional design of the survey does not allow these interpretations to be disentangled.

Regarding the dominance of attitudinal variables in the logistic regression model, B12 Likert-scale items on financial satisfaction, perceived constraint and goal setting behaviour show up among the top predictors by coefficient magnitude. Agreeing that financial situation limits goals is strongly associated with vulnerability, while reporting satisfaction with finance or strongly disagreeing that money is there to be spent both point toward resilience. These patterns align with Gathergood (2012), who documented the link between self-control, spending orientation and over-indebtedness in the United Kingdom. They also echo with a larger theme in behavioural economics: that attitudes and dispositions shape financial trajectories in ways that pure knowledge deficits do not fully explain.

The error analysis adds a further layer of nuance. The 19 vulnerable individuals whom XGBoost incorrectly classified as resilient are not, on the surface, typical vulnerable respondents. They save, they hold financial products, they report being satisfied with their finances. Many also expect to fund their retirement from their own savings, as the elevated B11_Savings_used variable indicates, even though their day-to-day position is too tight to absorb an immediate one-month shock. This profile is reminiscent of what Lusardi et al. (2011) described as the “financially fragile middle”: households that maintain the outward trappings of stability while lacking the liquid buffer to absorb even a modest shock. The fact that a sophisticated classifier struggles to identify these individuals suggests that their vulnerability is genuinely latent, visible only through the interaction of multiple subtle signals rather than any single red flag.

There is something genuinely surprising in this pattern that merits honest acknowledgment. This dissertation entered the analysis expecting objective knowledge to rank among the top five predictors, consistent with the weight it receives in the OECD/INFE framework and in much of the applied literature. That it did not, ranking below age, product holdings and even self-assessed knowledge, was not what the literature review had led the author to anticipate. Whether this reflects a genuine structural difference between the Portuguese context and the settings studied by Lusardi and colleagues, or simply the superior ability of tree-based models to capture behavioural interactions that logistic regression misses, remains an open question. Subsequent waves of the survey, or parallel analyses in other countries using the same OECD/INFE questionnaire, would be needed to assess the generalisability of this finding.

5.2 Policy Implications for Financial Education

The National Plan for Financial Education, coordinated by BdP, CMVM and ASF, identifies three strategic priorities for the 2021–2025 period: strengthening financial resilience, promoting digital financial literacy and contributing to sustainability (Conselho Nacional de Supervisores Financeiros, 2023a). The findings of this dissertation speak most directly to the first of these priorities.

If savings behaviour is the single most predictive feature of financial vulnerability, then interventions that help individuals establish and maintain a savings habit may yield larger resilience gains than programmes focused primarily on teaching financial concepts. This does not mean that knowledge is irrelevant; it is plausible that a baseline understanding of inflation, interest rates and risk is necessary for making informed savings decisions. But the evidence suggests that knowledge alone is insufficient, and that the translation from knowing to doing is where the greatest policy leverage lies.

Concretely, the results support three policy directions. The first concerns the design of savings nudges. Behavioural interventions such as automatic enrolment in savings schemes, default contributions to retirement plans and round-up programmes that channel spare change into savings accounts have shown promise in other national contexts. The role of B11 adds a complementary case for strengthening retirement-contribution defaults, given how many respondents expect to fund retirement from their own savings.

The second direction relates to targeting. The profile identified in Section 4.4 describes a population segment that would not be flagged by conventional vulnerability indicators based on income or employment status alone. These are mid-income individuals who save sporadically, hold some financial products and report moderate confidence in their financial plans, yet lack the buffer to weather a shock.

The third direction is attitudinal. The strong predictive signal carried by the B12 items suggests that attitudes toward money, particularly the tendency to view it as something to be spent rather than managed, are associated with vulnerability independently of income or knowledge. While changing deeply held attitudes is notoriously difficult, framing financial education around concrete goal-setting exercises, as opposed to abstract numeracy drills, may help shift dispositions in a direction that supports resilience. The survey's own finding that financial education is now a compulsory component of citizenship education in Portuguese schools provides a natural entry point for such an approach, provided the curriculum moves beyond textbook concepts toward applied decision-making skills.

5.3 Limitations

Several limitations should be borne in mind when interpreting these results. The most fundamental concerns the cross-sectional nature of the data. The survey captures a single snapshot of

respondents' financial situations, behaviours and attitudes at one point in time. It is therefore impossible to make causal claims about the direction of the relationships observed. Saving in the past year may protect against vulnerability, but it is equally possible that respondents who happen to be in a period of financial stability are simply more able to save. Longitudinal data, ideally from a panel that follows the same individuals across multiple survey waves, would be needed to establish temporal precedence and, by extension, to move closer to causal inference.

A second limitation relates to the target variable itself. Question B6 asks respondents to self-assess whether they could cover a one-month income shock. This is a hypothetical question, and the accuracy of respondents' self-reports is not independently verified. Some respondents may overestimate their capacity, particularly if they have never actually faced such a shock; others may underestimate it. The resulting measurement error is unlikely to be random, which could bias the model in ways that are difficult to quantify.

The sample size of 1,355 usable observations is modest by machine-learning standards. While the nested cross-validation framework provides unbiased performance estimates, the variance of those estimates across outer folds is non-trivial, as demonstrated by the standard deviations in Table 4.1. A larger sample would give tighter confidence intervals around the performance metrics and might also allow the detection of weaker predictive signals that are currently buried in noise.

The construction of the analytical sample introduces a further limitation. The exclusion of the 109 respondents who selected "Not applicable" to question B6, justified on definitional grounds in Section 3.2, has the side effect of narrowing the population to which the findings apply. Individuals without personal income, a group that plausibly contains a disproportionate share of students, home-makers and some pensioners, are absent from the modelling sample entirely. Their financial vulnerability is no less real, but it tends to operate at the household rather than the individual level, and it therefore falls outside the scope of an outcome defined in terms of one's own monthly income. The results should accordingly be read as describing individuals who report an income of their own, rather than the adult population as a whole.

The 19 "Do not know" responses were also excluded. Treating them as uninformative is an assumption rather than a certainty, since uncertainty about one's own resilience could itself signal precarity; given the small number of cases relative to the sample, however, the choice is unlikely to alter the substantive findings.

Finally, the feature importance rankings reported here are based on single-model-specific measures: impurity-based importance for random forest and coefficient magnitude for logistic regression. These measures have known limitations, including sensitivity to correlated features and scale dependence.

5.4 Future Research

The limitations identified above point directly to several avenues for future work. A panel extension of the survey, in which the same respondents are re-interviewed over time, would enable researchers to test whether changes in savings behaviour, attitude shifts or knowledge gains precede changes in vulnerability status, or vice versa. Such data would also permit the estimation of dynamic models that capture trajectories of resilience and fragility, rather than static classifications.

On the substantive side, replicating this analysis using financial literacy surveys from other countries that follow the OECD/INFE framework would test whether the dominance of behavioural over knowledge predictors is a general phenomenon or a feature specific to the Portuguese context. Cross-country comparisons could also illuminate the role of institutional factors, such as the availability of automatic savings mechanisms or the depth of financial product markets, in shaping the predictive landscape.

There is also an opportunity to integrate the survey data with administrative records, for example from the Portuguese tax authority or the social security system, to construct more objective measures of financial vulnerability. Such linkage raises privacy concerns that would need to be carefully managed, but it would address the self-report bias inherent in question B6 and potentially improve the precision of the predictive model.

The chapter that follows offers concluding reflections on what this analysis has achieved and what it means for the broader project of understanding and reducing financial vulnerability in Portugal.

6. Conclusion

This dissertation set out to answer a question with both analytical and practical dimensions: can financial vulnerability in the Portuguese population be reliably predicted from survey data, and what does the prediction tell us about the nature of vulnerability itself?

The answer to the first part is affirmative. Using data from the 4th Survey on the Financial Literacy of the Portuguese Population (2023), six supervised classification models were trained and evaluated through a rigorous nested cross-validation framework. XGBoost achieved the highest test-set ROC-AUC (0.913), followed closely by gradient boosting (0.909) and random forest (0.906). Even logistic regression, the simplest model in the battery, reached a test-set ROC-AUC of 0.897, indicating that the predictive signal in the data is strong enough to be captured by linear methods, though non-linear ensembles extract additional discriminative value.

The answer to the second part is more consequential and, in some respects, more challenging. The analysis reveals that financial vulnerability, defined as the inability to cover a one-month income shock without borrowing, is predicted far more strongly by what people do with their money than by what they know about financial concepts. The composite knowledge score, constructed from questions D2 through D7 covering numeracy, inflation, risk and digital financial literacy, ranks sixteenth out of 126 features in the random forest importance hierarchy. Savings behaviour, retirement planning confidence, household income and attitudinal self-assessments all outrank it substantially. This hierarchy does not deny the value of financial knowledge, but it does suggest that the pathway from knowledge to resilience is neither direct nor dominant. Behaviours and attitudes appear to mediate or even substitute for knowledge in ways that the standard financial literacy framework does not fully acknowledge.

The error analysis offers a practical complement to these aggregate findings. The individuals most likely to be misclassified as resilient by the best model are not the usual suspects: they are mid-income respondents who save sporadically, hold formal financial products and express satisfaction with their financial situation. Their vulnerability is hidden, residing not in an obvious income shortfall or knowledge gap but in a savings position too thin to absorb a sudden shock.

For the National Plan for Financial Education, the implications are twofold. Programmes that seek to improve financial resilience should not rely exclusively on knowledge transmission. Teaching people about compound interest or portfolio diversification has its place, but the evidence suggests that interventions aimed at establishing savings habits, building precautionary buffers and fostering goal-oriented financial planning are likely to yield more direct resilience gains. At the same time, the error profile highlights a segment of the population that existing risk indicators would miss, a group that merits attention precisely because it does not self-identify as vulnerable.

This dissertation does not pretend to have settled the relationship between financial literacy and vulnerability. Cross-sectional data cannot establish causation, self-reported outcomes carry measurement error, and the sample, while nationally representative, is modest in size for a machine-learning exercise. What it has done is demonstrate that a rich survey instrument, when subjected to modern classification techniques, can reveal patterns that descriptive analysis and standard regression approaches leave obscured. Whether those patterns prove robust across countries, time periods and alternative definitions of vulnerability is a question for the studies that follow this one.

References

- Abreu, M., & Mendes, V. (2010). Financial literacy and portfolio diversification. *Quantitative Finance*, 10(5), 515–528. <https://doi.org/10.1080/14697680902878105>
- Ampudia, M., van Vlokhoven, H., & Żochowski, D. (2016). Financial fragility of euro area households. *Journal of Financial Stability*, 27, 250–262. <https://doi.org/10.1016/j.jfs.2016.02.003>
- Anderloni, L., Bacchiocchi, E., & Vandone, D. (2012). Household financial vulnerability: An empirical analysis. *Research in Economics*, 66(3), 284–296. <https://doi.org/10.1016/j.rie.2012.03.001>
- Atkinson, A., & Messy, F.-A. (2012). *Measuring financial literacy: Results of the oecd/infe pilot study* (OECD Working Papers on Finance, Insurance and Private Pensions No. 15). OECD Publishing. Paris. <https://doi.org/10.1787/5k9csfs90fr4-en>
- Breiman, L. (2001). Random forests. *Machine Learning*, 45(1), 5–32. <https://doi.org/10.1023/A:1010933404324>
- Brunetti, M., Giarda, E., & Torricelli, C. (2016). Is financial fragility a matter of illiquidity? an appraisal for italian households. *Review of Income and Wealth*, 62(4), 628–649. <https://doi.org/10.1111/roiw.12189>
- Cawley, G. C., & Talbot, N. L. C. (2010). On over-fitting in model selection and subsequent selection bias in performance evaluation. *Journal of Machine Learning Research*, 11, 2079–2107.
- Chawla, N. V., Bowyer, K. W., Hall, L. O., & Kegelmeyer, W. P. (2002). SMOTE: Synthetic minority over-sampling technique. *Journal of Artificial Intelligence Research*, 16, 321–357. <https://doi.org/10.1613/jair.953>
- Chen, T., & Guestrin, C. (2016). XGBoost: A scalable tree boosting system. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 785–794. <https://doi.org/10.1145/2939672.2939785>
- Conselho Nacional de Supervisores Financeiros. (2023a). *Relatório do 4.º inquérito à literacia financeira da população portuguesa* (tech. rep.). Banco de Portugal, CMVM and ASF. Lisbon. <https://www.todoscontam.pt/>
- Conselho Nacional de Supervisores Financeiros. (2023b). *Survey on the financial literacy of the portuguese population 2023: Questionnaire and technical information* (tech. rep.). Banco de Portugal, CMVM and ASF. Lisbon. <https://www.todoscontam.pt/>
- Fernandes, D., Lynch, J., John G., & Netemeyer, R. G. (2014). Financial literacy, financial education, and downstream financial behaviors. *Management Science*, 60(8), 1861–1883. <https://doi.org/10.1287/mnsc.2013.1849>
- Friedman, J. H. (2001). Greedy function approximation: A gradient boosting machine. *The Annals of Statistics*, 29(5), 1189–1232. <https://doi.org/10.1214/aos/1013203451>

- Gathergood, J. (2012). Self-control, financial literacy and consumer over-indebtedness. *Journal of Economic Psychology*, 33(3), 590–602. <https://doi.org/10.1016/j.joep.2011.11.006>
- Hastie, T., Tibshirani, R., & Friedman, J. (2009). *The elements of statistical learning: Data mining, inference, and prediction* (2nd ed.). Springer. <https://doi.org/10.1007/978-0-387-84858-7>
- James, G., Witten, D., Hastie, T., & Tibshirani, R. (2021). *An introduction to statistical learning: With applications in R* (2nd ed.). Springer. <https://doi.org/10.1007/978-1-0716-1418-1>
- Khandani, A. E., Kim, A. J., & Lo, A. W. (2010). Consumer credit-risk models via machine-learning algorithms. *Journal of Banking and Finance*, 34(11), 2767–2787. <https://doi.org/10.1016/j.jbankfin.2010.06.001>
- Leo, M., Sharma, S., & Maddulety, K. (2019). Machine learning in banking risk management: A literature review. *Risks*, 7(1), 29. <https://doi.org/10.3390/risks7010029>
- Lusardi, A., & Mitchell, O. S. (2014). The economic importance of financial literacy: Theory and evidence. *Journal of Economic Literature*, 52(1), 5–44. <https://doi.org/10.1257/jel.52.1.5>
- Lusardi, A., Schneider, D. J., & Tufano, P. (2011). Financially fragile households: Evidence and implications. *Brookings Papers on Economic Activity*, 2011(1), 83–134. <https://doi.org/10.1353/eca.2011.0002>
- Moscatelli, M., Parlapiano, F., Narizzano, S., & Viggiano, G. (2020). Corporate default forecasting with machine learning. *Expert Systems with Applications*, 161, 113567. <https://doi.org/10.1016/j.eswa.2020.113567>
- OECD. (2022). *OECD/INFE toolkit for measuring financial literacy and financial inclusion 2022* (tech. rep.). OECD Publishing. Paris. <https://www.oecd.org/financial/education/2022-INFE-Toolkit-Measuring-Finlit-Financial-Inclusion.pdf>
- OECD. (2023). *OECD/INFE 2023 international survey of adult financial literacy* (tech. rep.). OECD Business and Finance Policy Papers, No. 39, OECD Publishing. Paris. <https://doi.org/10.1787/56003a32-en>
- Saito, T., & Rehmsmeier, M. (2015). The precision-recall plot is more informative than the ROC plot when evaluating binary classifiers on imbalanced datasets. *PLOS ONE*, 10(3), e0118432. <https://doi.org/10.1371/journal.pone.0118432>
- van Rooij, M., Lusardi, A., & Alessie, R. (2011). Financial literacy and stock market participation. *Journal of Financial Economics*, 101(2), 449–472. <https://doi.org/10.1016/j.jfineco.2011.03.006>
- Varma, S., & Simon, R. (2006). Bias in error estimation when using cross-validation for model selection. *BMC Bioinformatics*, 7(91). <https://doi.org/10.1186/1471-2105-7-91>

A. Variable Descriptions

Table A.1 provides a summary of the key variable groups used in the analysis. A full listing of all 229 raw survey variables can be found in the original questionnaire (Conselho Nacional de Supervisores Financeiros, 2023b).

Table A.1. Summary of variable groups used in the modelling pipeline, with survey section, key items and encoding approach.

Section	Key Variables	Description and Encoding
A: Profile	A1 (region), A2 (gender), A3 (age), A4 (education), A5 (employment), A6 (nationality), A7 (household composition)	Region and gender are treated as categorical and one-hot encoded. Age is numeric. Education and employment status are ordinal-encoded where a natural ordering exists; otherwise one-hot encoded. Household composition items (A7) are multcoded binary indicators.
B: Planning & Behaviour	B1–B5 (budgeting, savings methods), B6 (target variable), B10 (retirement confidence), B11 (retirement funding sources), B12 (attitudinal Likert items)	B4 and B5 are multcoded: each option becomes a binary column. B6 is the target (Yes=1, No=0; other responses excluded). B7 and B9 were removed to prevent data leakage. B10 is ordinal (0–5). B12 items are treated as categorical and one-hot encoded after RelativeFreqGrouper aggregation.
C: Financial Products	C1.1 (heard of), C1.2 (currently holds), C1.3 (purchased in last 2 years), C1.4 (most recent)	Each product in C1.1 and C1.2 yields a binary indicator. C1.4 is categorical (most recent product chosen) and one-hot encoded.
D: Knowledge & Income	D1 (self-assessment), D2–D7 (knowledge quiz items), D11 (income range)	D2–D7 are scored for correctness and summed into knowledge_score. D1 is ordinal (1–5). D10 items are binary. D11 is ordinal (1–4, with refusals treated as missing and dropped).

B. Confusion Matrices

Figure B.1 reproduces the confusion matrices for all six models on the held-out test set ($n = 271$). Each cell reports the absolute count of predictions. Row labels indicate the true class (Vulnerable or Resilient) and column labels indicate the predicted class.

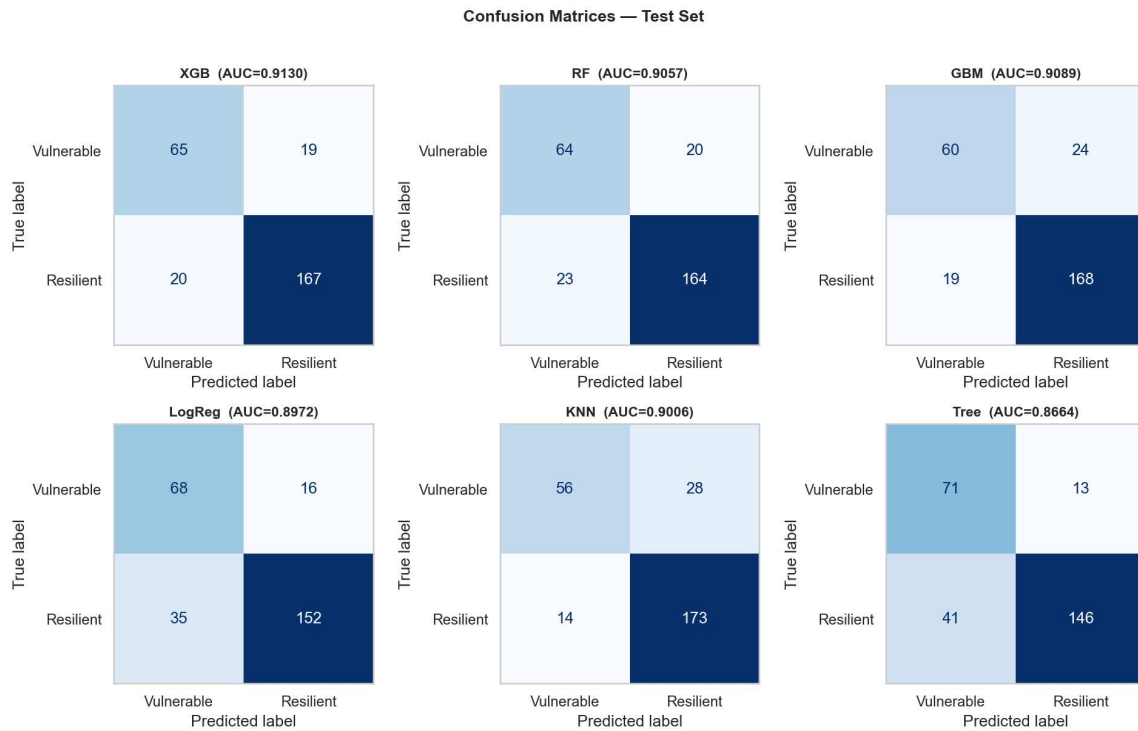


Figure B.1. Confusion matrices for all six models (test set, $n = 271$). XGBoost achieves the best balance between true positives (167) and true negatives (65), while logistic regression shows higher specificity at the cost of more false negatives.

C. Hyperparameter Search Spaces

Table C.1 reports the hyperparameter grids used in the inner loop of the nested cross-validation (RandomizedSearchCV, 30 iterations, 5-fold stratified CV, scored on ROC-AUC).

Table C.1. Hyperparameter search spaces for each model (RandomizedSearchCV, n_iter=30, 5-fold stratified CV, scoring='roc_auc'). Models whose total grid size was smaller than 30 were searched exhaustively.

Model	Parameter	Search Space
Logistic Regression	C	[0.001, 0.01, 0.1, 1, 10, 100]
	penalty	[l1]
	solver	liblinear (fixed)
	class_weight	[balanced]
Decision Tree	max_depth	[4, 5, 6, 8]
	min_samples_split	[4, 5, 6, 10]
	min_samples_leaf	[5, 6, 7, 10]
	class_weight	[balanced]
Random Forest	n_estimators	[100, 200, 300]
	max_depth	[3, 5, 7, None]
	min_samples_leaf	[10, 20, 50]
	max_features	[sqrt, log2]
	class_weight	balanced (fixed)
KNN	n_neighbors	[5, 9, 15, 21, 31]
	weights	[uniform, distance]
	p	[1, 2]
Gradient Boosting	n_estimators	[50, 100, 200]
	learning_rate	[0.01, 0.1, 0.2]
	max_depth	[3, 4, 5]
	(no class_weight; GradientBoostingClassifier does not expose this parameter)	
XGBoost	n_estimators	[100, 200, 300]

(continued on next page)

(continued from previous page)

Model	Parameter	Search Space
	learning_rate	[0.01, 0.05, 0.1, 0.2]
	max_depth	[3, 4, 5, 6]
	subsample	[0.7, 0.8, 1.0]
	colsample_bytree	[0.7, 0.8, 1.0]
	scale_pos_weight	0.451 (ratio of minority to majority class; fixed)