
Certificate in Credit Risk Analytics in Python

Credit Scoring Models

Acquisition Cost – the total expense incurred to obtain a new borrower, including marketing, underwriting, and processing fees. Example: A bank spends \$150 on advertising and \$50 on staff time for each approved loan, resulting in an acquisition cost of \$200 per customer. Understanding acquisition cost helps firms set pricing strategies and evaluate the profitability of credit products.

Alternative Data – non-traditional information sources used to supplement or replace conventional credit bureau data, such as utility payments, mobile phone usage, social media activity, and rental histories. Practical application: Machine-learning models ingest alternative data to improve score accuracy for thin-file applicants. Challenge: Ensuring data privacy compliance and mitigating bias introduced by unconventional variables.

Annual Percentage Rate (APR) – the yearly cost of borrowing expressed as a percentage, incorporating interest, fees, and compounding effects. Example: A loan with a 6% nominal rate and \$100 in fees on a \$1,000 principal may have an APR of approximately 7.2%. APR is a key output of credit scoring models when estimating borrower affordability.

Balance Sheet Ratios – financial metrics derived from a borrower's balance sheet, such as debt-to-equity, current ratio, and leverage ratio. Application: These ratios serve as predictive variables in logistic regression or tree-based models to assess default risk. Challenge: Ratios can be distorted by accounting policies, requiring careful preprocessing.

Bootstrap Aggregating (Bagging) – an ensemble technique that builds multiple models on bootstrapped subsets of data and averages their predictions to reduce variance. Example: Random Forests are a bagging method applied to decision trees for credit scoring. Bagging improves model stability, especially when dealing with noisy financial datasets.

Calibration – the process of adjusting predicted probabilities from a scoring model so they align with observed default frequencies. Technique: Platt scaling or isotonic regression are common calibration methods. Accurate calibration ensures that a score of 0.02 Truly reflects a 2% default probability, which is crucial for risk-based pricing.

Coefficient of Determination (R^2) – a statistical measure indicating the proportion of variance in the dependent variable explained by the model. In credit scoring, R^2 is less informative than classification metrics but can be used for regression-based loss-given-default models. Note: High R^2 does not guarantee good discrimination between good and bad borrowers.

Confusion Matrix – a tabular representation of classification outcomes: True positives (TP), false positives

(FP), true negatives (TN), and false negatives (FN). From this matrix, metrics such as accuracy, precision, recall, and F1-score are derived. Practical use: Evaluating a logistic-regression credit score on a validation set.

Cross-Validation – a resampling technique that partitions data into K folds, training the model on K-1 folds and testing on the remaining fold, rotating through all folds. Benefit: Provides robust performance estimates and guards against overfitting. In credit risk analytics, stratified K-fold cross-validation is preferred to preserve default rates across folds.

Cut-off Score – the threshold probability or score above which a loan application is approved. Determining the optimal cut-off involves balancing acceptance rates, expected loss, and profitability. Example: A bank may set a cut-off of 0.30, Meaning applicants with predicted default probability below 30% are approved.

Decision Tree – a hierarchical model that splits data based on feature thresholds to create leaf nodes representing predicted outcomes. Application: CART (Classification and Regression Trees) are widely used for interpretable credit scoring models. Challenge: Trees can overfit; pruning and limiting depth are essential safeguards.

Default Probability (PD) – the likelihood that a borrower will fail to meet contractual obligations within a specified time horizon, typically one year. Formula: $PD = \text{Number of defaults} / \text{Total number of borrowers in the cohort}$. PD estimates feed into capital allocation, pricing, and provisioning.

Delinquency – the condition of a loan being past due, commonly measured in days (e.G., 30-Day delinquency). Relevance: Delinquency status is an early indicator of credit deterioration and often used as a target variable for short-term scoring models.

Discriminatory Power – the ability of a credit scoring model to separate good borrowers from bad borrowers. Measured by the Gini coefficient or the Area Under the Receiver Operating Characteristic Curve (AUC-ROC). Higher discriminatory power implies more effective risk differentiation.

Distributional Shift – changes in the statistical properties of input features or target variables between training and deployment environments. Example: Economic downturns can alter default rates, causing a model trained on pre-crisis data to under-predict risk. Monitoring and periodic model retraining mitigate this issue.

Elastic Net – a regularization technique that combines L1 (Lasso) and L2 (Ridge) penalties to shrink coefficients and perform variable selection. Use case: In high-dimensional credit datasets, Elastic Net helps prevent overfitting while retaining important predictors.

Empirical Bayes – a statistical approach that leverages observed data to estimate prior distributions, often used in hierarchical modeling of default rates across segments. Application: Improves PD estimates for small loan portfolios by borrowing strength from larger, related groups.

Ensemble Modeling – the practice of combining multiple predictive models to produce a single, often more accurate, output. Techniques include bagging, boosting, and stacking. In credit scoring, ensembles can achieve higher AUC-ROC than any individual model.

Feature Engineering – the process of creating, transforming, and selecting variables that enhance model performance. Common techniques include binning continuous variables, generating interaction terms, and encoding categorical data. Challenge: Excessive feature creation can lead to multicollinearity and overfitting.

Feature Importance – a measure indicating how much each predictor contributes to the model's decisions. In tree-based models, importance can be derived from impurity reduction; in linear models, from absolute coefficient values. Understanding importance aids interpretability and regulatory compliance.

Fine-Tuning – the adjustment of hyperparameters (e.g., Learning rate, max depth) after initial model training to optimize performance. Grid search, random search, and Bayesian optimization are common fine-tuning strategies for credit scoring algorithms.

Gini Coefficient – a metric derived from the Lorenz curve that quantifies the inequality of a distribution. In credit scoring, the Gini is expressed as twice the area between the ROC curve and the diagonal, ranging from 0 (no discrimination) to 1 (perfect discrimination). It is frequently reported alongside AUC.

Gradient Boosting Machine (GBM) – an ensemble method that builds sequential decision trees, each correcting errors of its predecessor. Popular implementations include XGBoost, LightGBM, and CatBoost. GBMs often achieve state-of-the-art performance on credit scoring datasets but require careful regularization to avoid overfitting.

Imbalanced Data – a situation where the number of good borrowers vastly exceeds the number of defaults, common in credit portfolios. Mitigation: Techniques such as SMOTE oversampling, under-sampling, and cost-sensitive learning help balance the training process.

Information Value (IV) – a statistic that quantifies the predictive power of a categorical or binned variable. Calculated as $\sum (WOE_i \times (Distribution_good_i - Distribution_bad_i))$. IV > 0.3 is considered strong; IV

Interpretability – the degree to which a model's predictions can be understood by humans. Linear models and shallow decision trees are highly interpretable, while deep neural networks are less so. Regulatory frameworks often demand transparent scoring models, making interpretability a key design consideration.

K-Fold Cross-Validation – a specific cross-validation technique where the dataset is split into K equally sized folds. Each fold serves as a validation set once, and performance metrics are averaged across K runs. Commonly K=5 or K=10 in credit risk projects.

Logistic Regression – a statistical classification method that models the log-odds of default as a linear combination of predictors. Widely used in credit scoring for its simplicity, interpretability, and ease of calibration. Coefficients represent the change in log-odds per unit increase in the predictor.

Loss Given Default (LGD) – the proportion of exposure that is not recovered after a borrower defaults. Expressed as a percentage, $LGD = (Exposure - Recovery) / Exposure$. LGD models often employ regression techniques on collateral, industry, and macroeconomic variables.

Macro-Economic Variables – external factors such as unemployment rate, GDP growth, and interest rates that influence credit risk. Incorporating macro variables into scoring models improves predictive power during economic cycles. Challenge: Timely data acquisition and lag effects must be accounted for.

Margin of Error – the range within which a sample estimate is expected to fall relative to the true population parameter, typically expressed at a 95% confidence level. In credit scoring, margin of error informs the reliability of PD estimates derived from limited samples.

Model Drift – the gradual degradation of model performance over time due to changes in borrower behavior, product mix, or economic conditions. Continuous monitoring of key performance indicators (KPIs) like AUC and KS statistic helps detect drift early.

Multicollinearity – a condition where two or more predictors are highly correlated, inflating variance of coefficient estimates and reducing model stability. Detection methods include variance inflation factor (VIF) analysis. Remedies involve removing redundant variables or applying dimensionality reduction.

Neural Network – a set of algorithms inspired by biological neurons, capable of modeling complex nonlinear relationships. Deep learning architectures (e.G., Feed-forward, recurrent) are increasingly explored for credit scoring, especially when large, unstructured data sources are available. Challenge: Balancing predictive gains against interpretability and regulatory acceptance.

One-Hot Encoding – a technique for converting categorical variables into binary indicator columns, ensuring that models treat each category as a separate feature. Essential for algorithms that cannot handle non-numeric inputs, such as linear regression and many tree-based methods.

Out-of-Bag (OOB) Error – an internal estimate of model error for bagging ensembles, computed using data points not included in each bootstrap sample. OOB error provides a convenient validation metric without needing a separate hold-out set.

Partial Dependence Plot (PDP) – a visual tool that shows the marginal effect of a single predictor on the predicted outcome, averaged over the distribution of other features. PDPs aid interpretation of complex models like GBMs by illustrating how changes in a variable influence default probability.

Performance Monitoring – the ongoing assessment of a deployed credit scoring model using live data. Key metrics include AUC, KS statistic, population stability index (PSI), and calibration curves. Alerts trigger model review or retraining when thresholds are breached.

Probability of Default (PD) Curve – a graphical representation of predicted PDs across score bins, often overlaid with observed default rates to assess calibration. A well-calibrated curve will closely align predicted

and observed values.

Quantile Binning – a method of discretizing continuous variables by dividing them into equal-frequency intervals (e.g., Deciles). Binning reduces noise, facilitates Weight of Evidence (WoE) calculation, and improves model stability.

Receiver Operating Characteristic (ROC) Curve – a plot of true positive rate (sensitivity) versus false positive rate (1-specificity) across varying cut-off thresholds. The area under the ROC curve (AUC) quantifies overall discriminative ability. In credit scoring, AUC values above 0.70 are generally considered acceptable.

Regularization – a set of techniques (L1, L2, Elastic Net) that penalize large coefficient values to prevent overfitting. Regularization is especially important in high-dimensional credit datasets where the number of predictors may approach or exceed the number of observations.

Risk-Adjusted Return on Capital (RAROC) – a metric that compares expected profit to risk capital, adjusting for the probability of loss. $\text{RAROC} = (\text{Expected Return} - \text{Expected Loss}) / \text{Economic Capital}$. Credit scoring models feed PD and LGD estimates into RAROC calculations for portfolio optimization.

Sample Weighting – assigning different importance to observations during model training, often used to correct class imbalance or reflect the true exposure of each loan. In logistic regression, weights can be set equal to the loan amount to prioritize high-value accounts.

Segmentation – dividing a borrower population into homogeneous groups based on characteristics such as income, geography, or product type. Segmented models may capture distinct risk patterns and enable targeted pricing strategies.

Shapley Values – a game-theoretic method for attributing contribution of each feature to a model's prediction. Shapley values provide consistent, locally accurate explanations, making them valuable for interpreting complex models like gradient-boosted trees.

Stability Index (PSI) – a statistical measure that compares the distribution of a variable between training and validation (or production) datasets. $\text{PSI} > 0.25$ Typically signals a significant shift, prompting model review.

Stratified Sampling – a technique that ensures each class (e.g., Default vs. Non-default) is proportionally represented in training and test splits. This approach preserves the original default rate, leading to more reliable performance estimates.

Supervised Learning – a class of machine-learning algorithms that learn a mapping from input features to a known target variable (e.g., Default flag). Credit scoring is a classic supervised learning problem, with labels derived from historical loan outcomes.

Target Leakage – the inadvertent inclusion of information that would not be available at the time of prediction, such as post-loan performance metrics. Leakage inflates apparent model performance and must

be eliminated during data preparation.

Temporal Validation – evaluating a model on a hold-out period that chronologically follows the training window, mimicking real-world deployment. Temporal validation reveals how well the model generalizes to future data, accounting for macro-economic trends.

Threshold Optimization – the process of selecting a cut-off score that maximizes a business objective, such as expected profit or risk-adjusted return. Optimization may involve solving a simple profit equation or running a grid search over possible thresholds.

Tree-Based Models – algorithms that partition the feature space into rectangular regions using a hierarchy of splits. Examples include CART, Random Forest, and Gradient Boosting. Tree-based models handle nonlinearities and interactions automatically, making them popular for credit scoring.

Underwriting – the assessment process through which lenders evaluate borrower creditworthiness before extending credit. Automated underwriting systems rely on scoring models to make rapid, consistent decisions.

Validation Set – a subset of data reserved for tuning model hyperparameters and assessing performance before final testing. In credit risk projects, a separate validation set helps guard against overfitting to the training data.

Weight of Evidence (WoE) – a transformation of categorical or binned numeric variables into log-odds ratios: $WoE = \ln(\text{Distribution_good} / \text{Distribution_bad})$. WoE encoding preserves monotonic relationships and simplifies logistic regression coefficient interpretation.

Yield Curve – a graph showing the relationship between interest rates and maturity dates for debt instruments. While not a direct scoring input, yield curve movements affect borrower cost of capital and can be incorporated into macro-economic features.

Z-Score – a statistical measure representing the number of standard deviations a data point lies from the mean. In credit risk, Z-scores are used for outlier detection and for standardizing variables before modeling.

Zero-Inflated Models – statistical techniques that handle datasets with an excess of zero outcomes, such as loans with no defaults in a short observation window. Zero-inflated Poisson or negative binomial models can be employed when modeling count-based loss events.

Algorithmic Bias – systematic discrimination that arises when a model's predictions disproportionately affect protected groups (e.g., Based on gender or ethnicity). Mitigation strategies include fairness constraints, re-weighting, and careful variable selection.

Back-Testing – the retrospective evaluation of a scoring model by applying it to historical data and comparing predicted outcomes to actual results. Back-testing validates model assumptions and informs

adjustments before deployment.

Bootstrapping – a resampling method that creates multiple datasets by sampling with replacement from the original data. Used for estimating confidence intervals of model metrics and for generating OOB error estimates in ensemble methods.

Business Rule Engine – a rule-based system that applies deterministic criteria (e.G., Minimum income, maximum debt-to-income) before or alongside statistical scoring. Business rules provide a safety net for extreme cases where the statistical model may be unreliable.

Calibration Plot – a visual tool that compares predicted probability bins to observed default rates, highlighting over- or under-prediction. A perfectly calibrated model lies on the 45-degree line.

Data Imputation – the process of filling missing values in a dataset. Common techniques include mean/median substitution, k-nearest neighbors, and model-based imputation. Proper imputation prevents loss of valuable observations and reduces bias.

Environ-mental Stress Testing – scenario analysis that evaluates model performance under adverse macro-economic conditions (e.G., Recession, high unemployment). Stress testing ensures that scoring models remain robust during economic shocks.

Feature Selection – the identification of a subset of predictors that contribute most to model performance. Methods include recursive feature elimination, mutual information ranking, and regularization paths. Reducing feature count improves interpretability and reduces computational cost.

Gaussian Naïve Bayes – a probabilistic classifier assuming feature independence and Gaussian distribution of continuous variables. Though simplistic, it can serve as a baseline model for credit scoring, especially when data are limited.

Hyperparameter Tuning – the adjustment of algorithm-specific settings (e.G., Number of trees, learning rate) that are not learned from the data. Proper tuning can significantly enhance model accuracy and generalization.

In-Sample vs. Out-of-Sample – In-sample performance refers to metrics calculated on the data used to train the model, while out-of-sample performance uses unseen data. A large gap indicates overfitting; stable models exhibit similar metrics across both.

Jensen's Inequality – a mathematical principle stating that the transformation of an average is not equal to the average of transformed values for convex functions. In credit risk, this underlies the bias introduced when aggregating predicted probabilities without proper weighting.

KPI Dashboard – a visual interface displaying key performance indicators such as AUC, KS, PSI, and profit per loan. Dashboards enable risk managers to monitor model health in real time and trigger alerts when

thresholds are breached.

Log-Odds – the natural logarithm of the odds of default; central to logistic regression. Converting predicted probabilities to log-odds simplifies linear interpretation of coefficient effects.

Monte Carlo Simulation – a computational technique that generates a large number of random scenarios to assess the distribution of outcomes (e.G., Portfolio loss). Monte Carlo methods incorporate PD, LGD, and exposure at default (EAD) to estimate capital requirements.

Noise Ratio – the proportion of random error relative to the signal in a dataset. High noise reduces model predictability and may necessitate dimensionality reduction or more robust algorithms.

Outlier Detection – the identification of observations that deviate markedly from the majority of data. Techniques include Z-score thresholds, isolation forests, and robust Mahalanobis distance. Handling outliers prevents distortion of model coefficients.

Performance Attribution – the decomposition of portfolio results into components (e.G., Pricing, selection, risk) to understand the contribution of the scoring model to overall profitability.

Quantitative Risk Management – the systematic application of statistical and mathematical methods to identify, measure, and control credit risk. Credit scoring is a core quantitative tool within this discipline.

Random Forest – an ensemble of decision trees built on bootstrapped samples with random feature selection at each split. Random Forests provide strong predictive performance and built-in measures of feature importance, while reducing overfitting compared to single trees.

Sample Bias – distortion arising when the training data are not representative of the target population, often due to selection criteria or data collection methods. Correcting sample bias may involve re-weighting or augmenting the dataset.

Scorecard – a tabular representation that assigns points to each predictor based on its contribution to the overall score. Traditional scorecards translate logistic-regression coefficients into integer points, facilitating manual underwriting and regulatory reporting.

Segmentation-Specific Models – separate scoring models built for distinct borrower segments (e.G., SME vs. Consumer). Tailored models capture unique risk drivers and improve discrimination within each segment.

Threshold-Based Alerts – automated notifications triggered when a borrower's score crosses predefined risk levels, prompting review or intervention. Threshold alerts support proactive risk management.

Unsupervised Learning – algorithms that discover patterns without labeled outcomes, such as clustering borrowers by behavior. While not directly used for scoring, unsupervised techniques can inform feature engineering and segmentation.

Variance Inflation Factor (VIF) – a diagnostic metric that quantifies the increase in variance of a regression coefficient due to multicollinearity. $VIF > 10$ often signals problematic correlation requiring remedial action.

Weighted Average Cost of Capital (WACC) – the average rate a company is expected to pay to finance its assets. In credit risk, WACC informs the discount rate used in expected loss calculations and profitability analysis.

eXtreme Gradient Boosting (XGBoost) – a high-performance implementation of gradient boosting that includes regularization, parallel processing, and tree pruning. XGBoost is a popular choice for credit scoring competitions due to its speed and accuracy.

Yield-to-Maturity (YTM) – the total return anticipated on a bond if held until it matures. Though more relevant to fixed-income analysis, YTM trends can be incorporated as macro indicators influencing borrower creditworthiness.

Zero-Coupon Bond – a bond that pays no periodic interest and is sold at a discount to face value. Its pricing dynamics provide insight into long-term interest rate expectations, which can affect credit risk modeling.

Zero-Inflated Poisson (ZIP) Model – a statistical model for count data with excess zeros, combining a Poisson distribution with a separate binary process for zero inflation. In credit risk, ZIP models may be applied to count of missed payments before default.