

MACHINE LEARNING APPROACHES TO DEFAULT PREDICTION IN THE AUTOMOTIVE SECTOR

ABORDAGENS DE APRENDIZADO DE MÁQUINA PARA PREVISÃO DE INADIMPLÊNCIA NO SETOR AUTOMOTIVO

Alisson Ribeiro da Silva

Mestrando em Governança, Tecnologia e Inovação pela Universidade Católica de Brasília (Brasília/Brasil).
Especialista em Ciências de Dados e Inteligência Artificial pelo Centro Universitário Internacional UNINTER (Curitiba/Brasil).
E-mail: alisson1504@gmail.com

Paulo Fernando Marschner

Doutor em Administração pela Universidade Federal de Santa Maria (Santa Maria/Brasil).
Professor na Universidade Católica de Brasília (Brasília/Brasil).
E-mail: paulo.marschner@p.ucb.br

Eduardo Amadeu Dutra Moresi

Doutor em Ciências da Informação pela Universidade de Brasília (Brasília/Brasil).
Professor na Universidade Católica de Brasília (Brasília/Brasil).
E-mail: moresi@p.ucb.br

Recebido em: 19 de abril de 2026
Aprovado em: 2 de julho de 2026
Sistema de Avaliação: Double Blind Review
RGD | v. 23 | n. 1 | p. 29-56 | jan./jun. 2026
DOI: <https://doi.org/10.25112/rgd.v23i1.4572>

ABSTRACT:

The objective of this research was to conduct a bibliometric analysis of the international literature on the use of machine learning techniques for default prediction, with a focus on the automotive sector. A total of 1.026 documents published between 2010 and 2025 in the Scopus database were examined, aiming to identify thematic clusters, emerging trends, and methodological challenges. The results indicate a rapidly consolidating scientific field, characterized by international expansion, thematic diversification, and methodological maturity. The evidence is organized into four main axes: (i) the consolidation of machine learning as the central approach in credit risk modeling, gradually replacing traditional statistical methods; (ii) methodological expansion, with emphasis on neural networks, reinforcement learning, and ensemble models; (iii) growing appreciation for algorithmic interpretability and transparency; and (iv) integration of multiple data sources, combining financial, textual, and technological information. Moreover, the study identified a lack of systematic investigations into the application of machine learning models to default prediction specifically within the automotive sector, revealing a relevant gap in the literature. The implications of these findings include the construction of a structured overview of key trends, authors, and methodologies, providing theoretical and empirical insights to advance future research and improve predictive credit risk models.

Keywords: Machine learning. Default prediction. Automotive sector.

RESUMO:

O objetivo desta pesquisa foi realizar uma análise bibliométrica da literatura internacional sobre o uso de técnicas de aprendizado de máquina para previsão de inadimplência, com foco no setor automotivo. Um total de 1.026 documentos publicados entre 2010 e 2025 na base de dados Scopus foram examinados, visando identificar agrupamentos temáticos, tendências emergentes e desafios metodológicos. Os resultados indicam um campo científico em rápida consolidação, caracterizado por expansão internacional, diversificação temática e maturidade metodológica. As evidências estão organizadas em quatro eixos principais: (i) a consolidação do aprendizado de máquina como a abordagem central na modelagem de risco de crédito, substituindo gradualmente os métodos estatísticos tradicionais; (ii) expansão metodológica, com ênfase em redes neurais, aprendizado por reforço e modelos de conjunto; (iii) crescente valorização da interpretabilidade e transparência algorítmica; e (iv) integração de múltiplas fontes de dados, combinando informações financeiras, textuais e tecnológicas. Além disso, o estudo identificou uma carência de investigações sistemáticas sobre a aplicação de modelos de aprendizado de máquina à previsão de inadimplência especificamente no setor automotivo, revelando uma lacuna relevante na literatura. As implicações dessas descobertas incluem a construção de uma visão geral estruturada das principais tendências, autores e metodologias, fornecendo insights teóricos e empíricos para impulsionar pesquisas futuras e aprimorar modelos preditivos de risco de crédito.

Palavras-chave: Aprendizado de máquina. Previsão de inadimplência. Setor automotivo.

1 INTRODUCTION

The advancement of Machine Learning (ML) algorithms has profoundly transformed the way organizations manage credit risk, expanding predictive capabilities and enabling more robust data-driven decision-making processes. This transformation accompanies the evolution of integrated enterprise systems (Enterprise Resource Planning – ERP) and the increasing availability of external credit databases provided by specialized bureaus. In the field of default prediction modeling, traditionally supported by logistic regression models, there has been a shift toward algorithms capable of handling nonlinearity, class imbalance, and heterogeneous variables (BROWN; MUES, 2012; WANG *et al.*, 2011).

Default is one of the main challenges in credit risk management, with a strong impact on the financial sustainability of companies. In the automotive sector, this issue is intensified by the prevalence of installment sales and financing operations. The digitalization of business processes and the growing volume of transactional data have driven the use of ML techniques as an alternative to traditional statistical methods, allowing credit decisions to be based on historical patterns and multiple data sources.

In 2025, the Brazilian automotive sector has faced credit contraction and a rise in default rates, reaching 5.7% in the first quarter compared to 3.6% in the same period of 2024. The volume of credit for vehicle financing declined by 4.3%, totaling BRL 57.7 billion, reflecting high interest rates and greater selectivity among financial institutions (ANFAVEA, 2025). The Selic rate, at 15% per year, the highest level since 2006, has increased the cost of credit and limited the payment capacity of households and firms (REUTERS, 2025). This context reinforces the need for analytical tools capable of anticipating financial risks and supporting more accurate financing decisions.

The increase in default within this sector compromises cash flow and restricts credit throughout the automotive supply chain. Understanding the determinants of payment default is therefore essential for developing mitigation strategies and improving credit assessment mechanisms. Conceptually, default corresponds to the failure to meet financial obligations within the contractual period (SILVA; NEVES, 2020), while credit risk represents the probability that a counterparty will fail to honor its obligations, generating financial losses (SAUNDERS; ALLEN, 2021). Both concepts highlight the need for accurate measurement and forecasting to ensure financial stability and operational efficiency.

Although the use of ML techniques in credit risk modeling has been expanding, little is known about how this movement manifests itself specifically within the automotive sector. Scientific production at this intersection appears to be dispersed across different areas, including finance, data science, and production engineering, which makes it difficult to understand its intellectual structure and evolutionary trajectory. This fragmentation also limits the identification of the extent to which advances in ML-based

default prediction have been incorporated into automotive finance research and practice. In light of this gap, the present study seeks to answer the following research question: How has the international scientific literature evolved regarding the application of machine learning techniques to default prediction, and to what extent has this knowledge been applied to the automotive sector? To address this question, this study conducts a bibliometric analysis of the international literature on the use of ML techniques for default prediction, with particular attention to their potential application in the automotive sector. A total of 1,026 documents published between January 2010 and September 2025 were analyzed from the Scopus database to identify thematic clusters, emerging trends, and methodological challenges.

The main findings include: (i) the consolidation of ML as the central framework for default prediction, gradually replacing traditional statistical models with more flexible and self-adaptive approaches; (ii) the expansion of methodological scope, with emphasis on neural networks, reinforcement learning, and ensemble models that enhance accuracy and generalization capacity; (iii) the growing emphasis on interpretability and algorithmic transparency, reflecting increasing concern with traceability and trust in applied models; and (iv) the integration of multiple data sources, combining financial, textual, and technological information to improve predictive robustness and broaden applicability across different contexts, such as small businesses and supply chains. Additionally, the study identified a lack of systematic investigations into the application of ML models specifically to default prediction in the automotive sector, revealing a relevant gap in the literature and an opportunity for future studies in this domain.

These findings contribute, both theoretically and empirically, to advancing the understanding of credit risk modeling by demonstrating the convergence of the literature around hybrid, explainable, and fair frameworks that seek to balance accuracy, adaptability, and equity in credit decisions. From a theoretical perspective, this convergence represents an important epistemological transformation of the field, which now integrates technical and ethical dimensions into the design of predictive models. From an empirical standpoint, there is a consolidation of practices guided by transparency, interpretability, and algorithmic accountability, reflecting the growing scientific and practical maturity of ML applications in financial risk management.

2 LITERATURE REVIEW

The study of default and credit risk originated in efforts to measure and predict the behavior of borrowers, aiming to reduce the uncertainty associated with financial decisions. Since the earliest econometric formulations, the central concern has been to develop quantitative tools capable of

supporting the judgment of analysts and financial institutions. Altman's (1968) Z-Score model represents a historical milestone in this process by combining multiple accounting indicators into a discriminant function capable of estimating the probability of bankruptcy. Altman's model inaugurated an evidence-based logic of assessment and objective risk measurement, allowing credit analysis to emerge as an autonomous scientific field within corporate finance.

In the following decades, **credit risk assessment practices evolved** with the emergence of credit scoring models, which expanded the use of demographic and behavioral variables to estimate the individual risk of each borrower. These models, originally grounded in linear regressions and classical statistical techniques, became a cornerstone of banking and consumer credit policies. According to Thomas, Crook, and Edelman (2017), the consolidation of credit scoring represented not only a technical advancement but also an epistemological transition from approaches based on intuition and subjective judgment to methods grounded in statistical evidence and probabilistic reasoning. The transition toward evidence-based credit assessment defined the identity of the field as an applied branch of statistical modeling focused on balancing predictive accuracy and interpretability, establishing the foundations upon which contemporary ML-based approaches to credit risk would later develop.

With the advancement of computational power and data availability, default predictive modeling progressively incorporated ML techniques, overcoming the limitations of traditional linear models (BROWN; MUES, 2012). The adoption of ML techniques reflects a methodological transition in which models learn complex patterns from imbalanced datasets and capture nonlinear interactions among financial, demographic, and behavioral variables, thereby enhancing the predictive performance of default models.

Lessmann *et al.* (2015) conducted a comprehensive benchmarking study comparing the performance of various classification algorithms applied to credit scoring, establishing a major methodological benchmark in the field. The authors demonstrated that ML-based models, such as Random Forest and Gradient Boosting, significantly outperform traditional logistic regression in terms of accuracy and robustness, particularly in complex datasets. Their findings stimulated a new research agenda focused on systematic experimentation and comparative evaluation of classification methods in real-world credit contexts.

Recent literature highlights the growing use of algorithms such as Random Forest, Gradient Boosting, and Artificial Neural Networks (Xiao *et al.*, 2025). Although logistic regression remains widely employed due to its interpretability and simplicity, ensemble learning methods have shown superior performance in large and complex datasets.

The integration of structured and unstructured data has enhanced the **predictive capability** of credit risk models. Wu *et al.* (2025) explored textual variables in credit reports, while Zhang *et al.* (2021) and Liu *et al.* (2024) emphasized the potential of domain adaptation and feature engineering to tailor models to specific corporate contexts. Collectively, these developments indicate a movement toward integrating quantitative and heuristic approaches, aiming to combine predictive performance with greater model interpretability and operational applicability. At the same time, the increasing diversity of methods, data sources, and analytical approaches has contributed to the expansion of a multidisciplinary body of knowledge spanning finance, data science, and management research.

Despite this progress, a relevant scientific gap persists regarding **the integration** of internal corporate data (such as ERP records, payment histories, and credit limits) with external indicators, including bureau scores and macroeconomic variables. This gap is particularly evident in studies applied to the automotive sector, whose operational structure combines credit, logistics, and after-sales services, requiring models adapted to the complexity of business operations. The persistence of these integration challenges suggests that the evolution of credit risk modeling cannot be fully understood solely through methodological advances, but also through the contexts in which such models are applied.

Previous studies have emphasized that credit risk environments differ substantially across economic sectors due to variations in financing structures, operational characteristics, and exposure profiles (Saunders & Allen, 2021; Breeden, 2021). The automotive sector presents particular characteristics that distinguish it from traditional consumer and banking credit environments. **Lin et al. (2025) highlight the growing relevance of credit risk assessment in the automotive sector, arguing that the expansion of automobile financing has increased the importance of effective risk assessment mechanisms.** Default prediction in this segment involves distinct characteristics related to loan conditions, borrower profiles, vehicle-related factors, and repayment capacity. Unlike traditional credit environments, automobile credit requires models capable of capturing complex interactions among financial, behavioral, and operational variables. Using data from an automotive financing company, the authors demonstrated that models such as XGBoost combined with SHapley Additive exPlanations (SHAP) can improve predictive performance while preserving interpretability. Their results underscored the importance of borrower- and loan-related characteristics in assessing automotive credit risk, reinforcing the need for predictive models that are explainable and adapted to the specificities of the sector.

Despite the economic relevance of the automotive industry, the application of machine learning techniques to default prediction in this context remains fragmented and underexplored in the scientific literature. Most studies continue to focus on banking, consumer lending, or fintech environments, leaving important opportunities for sector-specific investigations. The fragmented nature of the literature

becomes even more evident when considering the growing diversification of themes associated with credit risk modeling, including explainability, organizational decision-making, and other emerging discussions related to the responsible use of ML. These emerging themes reflect a broader shift in the field from a narrow focus on predictive accuracy toward a more comprehensive view of responsible and transparent decision-making.

Recent advances in credit risk modeling have also increased interest in the interpretability of machine learning models. Although complex algorithms often achieve superior predictive performance, their decision-making processes may be difficult to understand, limiting their practical adoption in organizational contexts. Consequently, explainable artificial intelligence (XAI) approaches have received growing attention as mechanisms for improving transparency and supporting managerial decision-making. In addition, concerns regarding potential algorithmic biases and fairness have emerged as complementary challenges, particularly in contexts where automated credit decisions may affect access to financial resources and business opportunities.

Beyond predictive performance, the growing adoption of machine learning in credit risk assessment has introduced broader organizational challenges related to governance, transparency, and the effective integration of analytical models into decision-making processes. As argued by Zancan, Passador and Passador (2023), AI-based systems should be understood not only as analytical tools but also as mechanisms that support managerial decision-making in complex organizational environments. This perspective suggests that the evolution of credit risk modeling increasingly involves balancing predictive capability, interpretability, and organizational applicability.

Taken together, these developments reveal a rapidly expanding and increasingly interdisciplinary field characterized by methodological diversification, sector-specific applications, and the growing integration of machine learning techniques into different organizational contexts. Despite this evolution, the literature remains fragmented across different research domains, making it difficult to obtain a comprehensive understanding of its intellectual structure, thematic evolution, and emerging research fronts. Therefore, a bibliometric investigation is warranted to systematically map **the intellectual development of this field**, identify dominant themes and research gaps, and explore the extent to which advances in machine learning-based default prediction have been incorporated into the automotive sector.

3 METHODOLOGY

This research adopts an integrated approach combining quantitative and qualitative analysis, suitable for systematic reviews and bibliometric studies. This combination integrates statistical rigor with contextual interpretation, allowing the mapping of scientific trends and the identification of theoretical gaps in the field of default prediction (MORESI; PINHO; COSTA, 2021). Data collection was conducted using the Scopus database, between August and September 2025, based on the search expression presented in Table 1. The strategy encompassed three thematic axes: (i) default and credit risk, (ii) ML techniques, and (iii) business and automotive contexts.

Documents published between 2010 and 2025 were considered, in either English or Portuguese, including articles, reviews, and conference papers. Metadata were refined to ensure consistency and reproducibility, comprising: (i) removal of duplicates and incomplete records; (ii) standardization of author and institutional names; and (iii) the development of a thesaurus to unify synonyms and linguistic variations (for example, credit risk and credit-risk assessment). The process followed the methodological recommendations of Zupic and Čater (2015).

Table 1 – Initial search expression.

(TITLE-ABS-KEY (("payment delinquen*" OR "credit risk*" OR "commercial credit" OR "customer default*" OR "credit analysis" OR "credit default" OR "loan default" OR "default prediction" OR "non-performing loan*" OR "creditworthiness" OR "credit scoring" OR "delinquency prediction" OR "trade credit" OR "corporate credit risk") AND ("machine learning" OR "artificial intelligence" OR "predictive model*" OR "data mining" OR "neural network*" OR "deep learning" OR "random forest" OR "support vector machine*" OR "gradient boosting" OR "xgboost" OR "ensemble learning" OR "logistic regression" OR "explainable AI" OR "feature selection" OR "classification algorithm*" OR "Knowledge Discovery in Databases" OR "KDD") AND ("retail" OR "consumer" OR "B2B" OR "B2C" OR "wholesale" OR "business" OR "SME*" OR "dealer" OR "supplier" OR "service industry" OR "commercial credit" OR "trade receivables" OR "accounts receivable" OR "supply chain finance" OR "ERP system" OR "enterprise resource planning" OR automotive OR "automotive aftermarket" OR "auto parts" OR "spare parts" OR "tire*" OR "tyre*" OR "vehicle maintenance" OR "dealer*" OR "distributor*" OR "fleet management" OR "trucking" OR "road freight" OR "logistics" OR "sale*")) AND NOT TITLE-ABS-KEY ("bank*" OR "credit card" OR "mortgage")) AND PUBYEAR > 2009 AND PUBYEAR < 2026 AND (LIMIT-TO (LANGUAGE , "English") OR LIMIT-TO (LANGUAGE , "Portuguese"))

Source: prepared by the authors.

The search string was designed to ensure both precision and comprehensiveness in retrieving relevant studies. The strategy combines three thematic blocks connected by Boolean operators: (i) credit risk and default (payment delinquency, credit risk, loan default); (ii) artificial intelligence methods and predictive modeling (ML, neural networks, random forest, XGBoost, explainable AI); and (iii) business and automotive contexts (retail, automotive, ERP systems, dealers, logistics). Studies focused on banking and personal credit were excluded through the NOT operator in order to reduce the dominance of traditional banking studies and ensure the relevance of the retrieved results to sectoral and commercial credit contexts.

The quantitative analysis was conducted using three complementary bibliometric tools: VOSviewer (v.1.6.20), Bibliometrix/Biblioshiny (RStudio, v.4.3.3), and Gephi (v.0.10.1). In VOSviewer, the Leiden algorithm was applied (resolution = 1.0; minimum of five keyword occurrences), with normalization based on association strength and clustering according to eigenvector centrality (CALLON *et al.*, 1983). Bibliometrix was employed to generate classical bibliometric indicators, including the number of publications per year, growth rate, international collaboration, most productive sources, average citations, and the Three-Field Plot (authors–institutions–countries). Finally, Gephi was used to compute network metrics (degree, betweenness centrality, and modularity class), allowing the identification of communities and co-authorship and co-citation relationships. The combined use of these tools ensures complementarity: VOSviewer provides semantic visualization, Bibliometrix supplies descriptive statistics, and Gephi reveals the relational structure of the networks.

Based on the quantitative results, a qualitative analysis of the identified clusters was then conducted. This stage followed the principle of methodological triangulation (Denzin, 1978), linking centrality and frequency metrics to the contextual reading of the most cited and most recent studies. This procedure ensures an integrated interpretation of emerging trends and contributes to identifying theoretical gaps and opportunities for future research.

4 RESULTS

The metadata was exported in CSV format and processed in Bibliometrix (RStudio) to calculate the main bibliometric indicators. Table 2 presents a synthetic overview of the analyzed dataset, providing an overview of the distribution of documents, sources, authors, and impact metrics within the corpus.

Table 2 – General indicators of the bibliographic sample (2010–2025).

Indicator	Value	Interpretation
Analysis period	2010–2025	15 years of scientific production
Total documents	1,026	Moderate and cohesive research corpus
Sources (journals, conferences, etc.)	377	High diversity of publication channels
Annual growth rate	12.81%	Continuous expansion of academic interest
Average document age	4.48 years	Recent and dynamic production
Average citations per document	16.66	Good impact and visibility
Cited references	6,515	Dense and interconnected bibliographic base
Total authors	1,940	Expressive scientific community
Average co-authors per article	6.01	Strong interdisciplinary collaboration
International collaboration	17.64%	Moderate level of global exchange

Source: prepared by the authors.

The main publication areas were Computer Science (42%), Business, Management & Accounting (28%), and Engineering (18%), highlighting the multidisciplinary nature of the field. The predominance of journal articles and conference papers reflects the dynamism of the research area and the rapid dissemination of new findings, while the presence of systematic reviews indicates a certain level of theoretical maturity.

The combined use of VOSviewer and Gephi enabled the identification of keywords with the highest eigenvector centrality (CALLON *et al.*, 1983) and the mapping of co-occurrence networks among core concepts such as credit risk modeling, default prediction, data integration, and explainable AI. These networks reveal the conceptual nuclei and structural connections that underpin the literature on default prediction. A total of 3,437 Keywords Plus and 2,134 author-provided keywords were identified, reflecting a consolidated vocabulary, although with variations among indexers and researchers. This difference suggests conceptual nuances and potential emerging trends to be further explored in co-occurrence analyses.

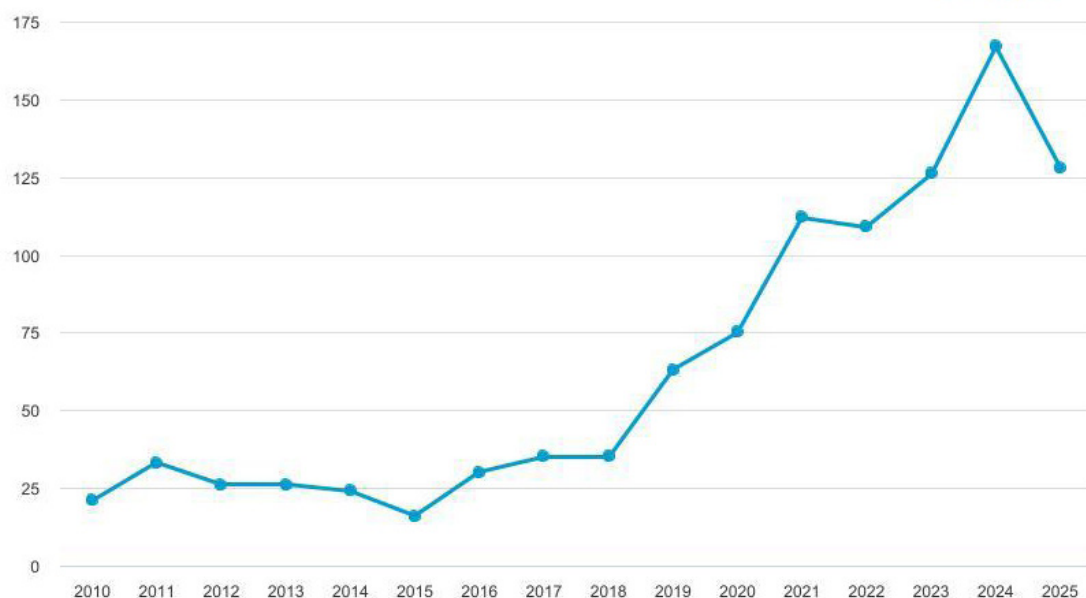
Finally, the collaborative profile of the scientific community stands out, characterized by the absence of single-authored publications and strong institutional collaboration. This pattern indicates the interdisciplinary and applied nature of the field, which combines computational techniques with business applications.

The word cloud (Figure 1) was generated from the keywords of the analyzed articles, visualizing the most frequent and relevant terms in literature. The size of each term reflects its frequency and relative importance within the corpus, highlighting core concepts such as ML, credit scoring, credit risk, logistic regression, and default prediction, which constitute the thematic nucleus of research on default prediction using ML algorithms. Additional terms such as ensemble learning, deep learning, feature selection, and explainable artificial intelligence illustrate the contemporary trend toward the use of hybrid and interpretable models in credit risk analysis. Thus, the word cloud provides an intuitive and exploratory representation of the conceptual structure of the field, allowing the identification of key research focuses and emerging methodological directions.

[illegible]

Figure 2 shows the evolution of publications between 2010 and 2025. A gradual increase can be observed until 2018, followed by a sharp expansion starting in 2019, reflecting the consolidation of academic interest in the use of ML and artificial intelligence techniques applied to default prediction and credit risk assessment. The peak occurs in 2024, with 167 publications, representing the highest level of scientific output on the topic, possibly driven by advances in predictive modeling and the widespread adoption of algorithms such as XGBoost, deep neural networks, and explainable learning approaches. In 2025, a slight decline to 128 studies is observed, which may reflect a natural adjustment in annual production, the consolidation of the field, or the fact that the year is still ongoing, which may limit the number of indexed publications. Overall, the trend indicates a process of thematic maturation, particularly after 2020, when the integration of business data, predictive algorithms, and risk management began to occupy a central position in scientific literature.

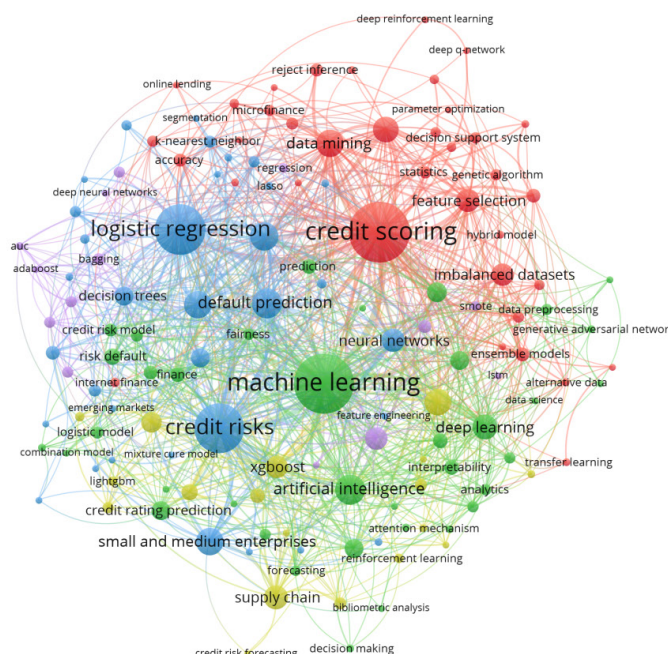
Figure 2 – Publication Trends.



Source: prepared by the authors.

The VOSviewer software was used to construct the keyword co-occurrence network (Figure 3). This visualization makes it possible to identify thematic relationships as well as the intensity of connections among the most recurrent terms, revealing interconnected clusters that reflect the main research axes. To ensure consistency and accuracy of the results, a thesaurus was developed to standardize synonyms, eliminate duplicates, and unify linguistic variations, for example, treating "credit risk," "credit risks," and "credit-risk assessment" as equivalent terms. This standardization prevents thematic fragmentation, reduces distortions in the interpretation of conceptual centrality, and ensures the reliability of co-occurrence networks, thereby enhancing comparability across studies and the robustness of analyses produced by VOSviewer and complementary tools such as Gephi and Bibliometrix.

Figure 3 – Keyword Co-occurrence Network.



Source: prepared by the authors.

The red cluster concentrates terms related to predictive modeling and data preprocessing, such as credit scoring, data mining, feature selection, and imbalanced datasets, indicating a recurring concern with performance and data quality. The blue cluster highlights the classical methodological axis, centered on logistic regression, decision trees, and default prediction, representing the statistical foundations still widely employed in credit risk modeling. The green cluster encompasses advances in ML and deep learning, with emphasis on artificial intelligence, ensemble models, and interpretability, signaling an evolution toward more sophisticated and explainable methods. The yellow cluster is associated with corporate applications, linking terms such as supply chain, small and medium enterprises, and credit risk forecasting, which reflects the growing interest in expanding these approaches to business and operational contexts. Overall, the network reveals high density and interconnection among terms, emphasizing the interdisciplinary nature of the field and the progressive integration of techniques, applications, and theoretical foundations in research on default prediction and credit risk.

From this mapped conceptual structure, it becomes possible to interpret not only consolidated thematic cores but also the absences and gaps within the network, which point to promising opportunities for investigation. A detailed examination of this configuration shows that the literature on

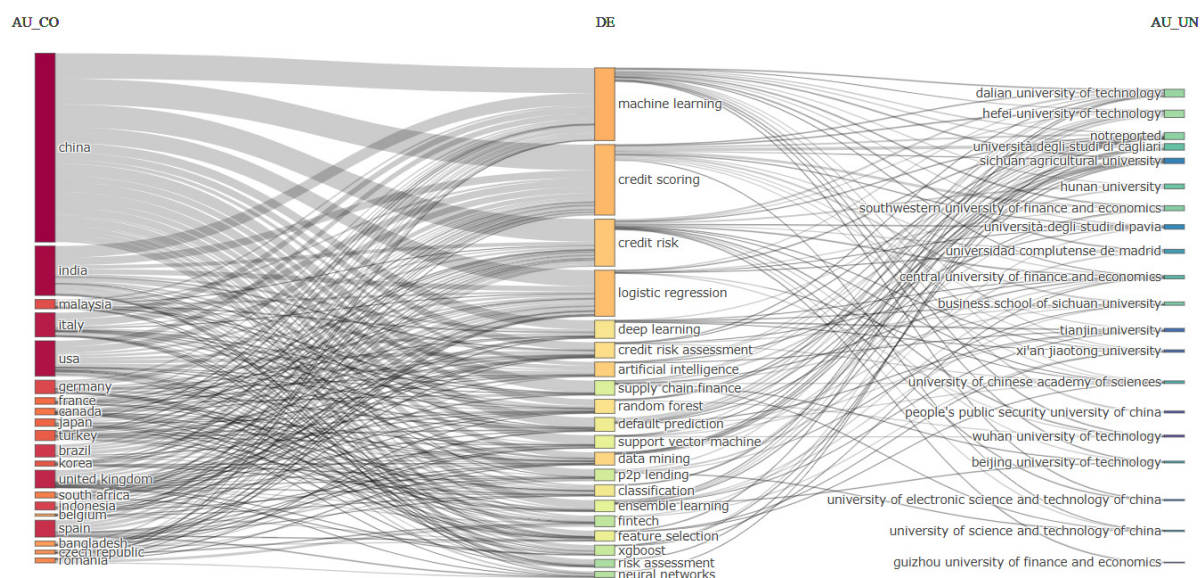
default prediction and credit risk is concentrated mainly in banking, credit card, and consumer financing contexts, with no explicit occurrence of terms or co-occurrences related to the automotive sector. Although this sector occasionally appears as an implicit field of application, it has not yet been addressed autonomously, remaining diluted within broader studies on supply chains or commercial credit. This pattern reveals a predominance of research focused on financial institutions, with limited attention to interfirm credit mechanisms, whose risk structure differs substantially from traditional banking credit. Thus, the bibliometric analysis supports the identification of a relevant gap in literature: the absence of systematic investigations on the application of ML models to default prediction in the automotive sector.

Complementing the analysis of the keyword co-occurrence network, the bibliometric investigation broadens the understanding of the conceptual and institutional structure of the field “ML Applied to Default Prediction.” The results indicate a rapidly expanding scientific domain, characterized by high collaboration, growing relevance, and moderate yet significant impact. This configuration confirms that it is a young and dynamic area, still in the process of consolidation, but already exhibiting consistent signs of maturation and international recognition.

The integrated analysis of relationships among countries, institutions, and keywords, represented in the Three-Field Plot (Figure 4), reinforces this perception. The diagram illustrates how specific terms connect to particular geographic and institutional contexts, making it possible to identify production hubs and thematic specialization centers. The scientific output is concentrated primarily in China, followed by India and the United States, with relevant contributions from Spain, the United Kingdom, Indonesia, Italy, France, South Korea, Japan, Brazil, Canada, and Germany, a scenario that reveals Asian predominance, though with significant European participation.

In the keyword axis, ML and credit scoring stand out, followed by terms such as credit risk, logistic regression, deep learning, artificial intelligence, supply chain finance, random forest, data mining, and fintech, indicating an emphasis on ML and data analysis applications for financial risk management. Regarding institutions, the Dalian University of Technology, the Hefei University of Technology, and the Southwestern University of Finance and Economics are prominent, with international collaborations involving universities such as Università degli Studi di Cagliari (Italy) and Universidad Complutense de Madrid (Spain). The category “not reported” refers to affiliations that are missing or not standardized

Figure 4 - Relationships between countries, universities, and keywords.



Source: prepared by the authors.

The predominant connection between China and the terms ML and credit scoring highlights the thematic emphasis of Chinese researchers and institutions, reinforced by links with several universities in the country. The dispersion of other terms across different countries and institutions indicates global interest and possible institutional specializations, such as the strong association with supply chain finance, which suggests established research groups in this area. This overview makes it possible to map key actors (countries, universities, and authors), monitor trends, identify potential collaborators, and position the research within the international landscape, in addition to informing policies and funding strategies that explain geographical concentrations.

The analysis of the most cited articles provides insight into the scientific impact and structural influence of studies in the field of ML applied to default prediction. The predominance of the journal *Expert Systems with Applications* confirms its role as the main outlet for disseminating research in this domain, evidencing a consolidated community and a scope focused on intelligent systems and analytical applications. Among the most prominent studies, those with high local citation (LC) and global citation (GC) indices constitute the core of influence both within the thematic scope and the international landscape. The article by Brown and Mues (2012), with LC = 34 and GC = 557, represents a seminal reference, while Dumitrescu *et al.* (2022) (LC = 13; GC = 269) consolidates the application of nonlinear methods to credit modeling.

Conversely, studies with high LC and low GC, such as Bequé and Lessmann (2017), stand out for their intrinsic relevance within the field, although they have not yet achieved wide international dissemination, which may be associated with a narrower thematic focus or shorter circulation time. In contrast, studies with low LC and high GC, such as Wang *et al.* (2012), demonstrate significant global impact but lower alignment with the specific scope of this analysis. The normalized metrics (NLC and NGC) help to account for temporal differences and to identify emerging trends. In this regard, Belhadi *et al.* (2025) presents strong performance in both metrics, signaling growing potential for impact, albeit recent.

Table 3 - Most cited documents.

	Authors	Title	Journal	LC	GC
1	Brown, I.; Mues, C.	<i>An experimental comparison of classification algorithms for imbalanced credit scoring data sets</i>	Expert Systems with Applications	34	557
2	Bequé, A.; Lessmann, S.	<i>Extreme Learning Machines for Credit Scoring: An Empirical Evaluation</i>	Expert Systems with Applications	16	127
3	Dumitrescu, E.; Hué, S.; Hurlin, C.; Tokpavi, S.	<i>Machine learning for credit scoring: Improving logistic regression with non-linear decision-tree effects</i>	European Journal of Operational Research	13	269
4	Bahnsen, A. C.; Aouada, D.; Ottersten, B.	<i>Example-Dependent Cost-Sensitive Logistic Regression for Credit Scoring</i>	ICMLA Proceedings	12	124
5	Chang, Y.-C.; Chang, K.-H.; Wu, G.-J.	<i>Application of eXtreme Gradient Boosting Trees in the Construction of Credit Risk Assessment Models for Financial Institutions</i>	Applied Soft Computing	11	250
6	Blanco, A.; Pino-Mejías, R.; Lara-Rubio, J.; Rayo, S.	<i>Credit Scoring Models for the Microfinance Industry Using Neural Networks: Evidence from Peru</i>	Expert Systems with Applications	10	153
7	Bücker, M.; Szepannek, G.; Gosiewska, A.; Biecek, P.	<i>Transparency, Auditability, and Explainability of Machine Learning Models in Credit Scoring</i>	Journal of the Operational Research Society	9	106
8	Ariza-Garzón, M. J.; Arroyo, J.; Caparrini, A.; Segovia-Vargas, M. J.	<i>Explainability of a Machine Learning Granting Scoring Model in Peer-to-Peer Lending</i>	IEEE Access	7	113
9	Chang, Y.-C.; Chang, K.-H.; Chu, H.-H.; Tong, L.-I.	<i>Decision Tree-Based Short-Term Default Credit Risk Assessment Models</i>	Communications in Statistics – Theory and Methods	7	26
10	Belhadi, A.; Kamble, S. S.; Mani, V.; Benkhathi, I.; Touriki, F. E.	<i>An Ensemble Machine Learning Approach for Forecasting Credit Risk of Agricultural SMEs' Investments in Agriculture 4.0 through Supply Chain Finance</i>	Annals of Operations Research	5	50

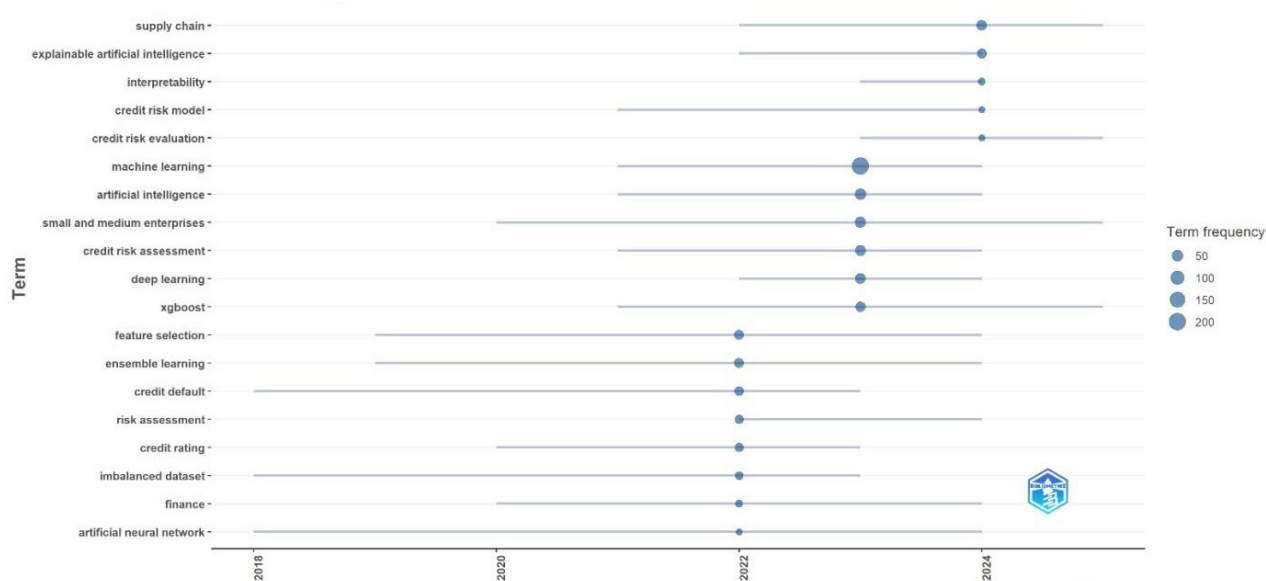
Source: prepared by the authors.

Overall, the analysis of the most cited articles goes beyond simple reference counts. It reveals the intellectual and thematic vectors that structure the field, highlighting contributions with high theoretical and applied impact. Understanding why certain studies have become central, whether due to methodological

innovation, practical applicability, or alignment with strategic topics, provides valuable insights for positioning new research and strengthening scientific integration within the international ecosystem.

The understanding of the conceptual structure becomes deeper when the intellectual dimension of the field is also examined, revealing its theoretical foundations and the evolutionary trajectories of its main themes. This dimension shows not only the consolidation of statistical and computational approaches but also the gradual incorporation of artificial intelligence techniques into credit risk management, a movement that, although still concentrated in the traditional financial context, is beginning to expand into specific sectoral applications, such as the automotive industry. In this regard, the analysis of *Trend Topics* complements and enriches this interpretation by highlighting the temporal dynamics of dominant themes in *credit scoring* and related areas. The chart (Figure 5) shows the evolution of research topics and the consolidation of thematic axes, in which the size of the bubbles and the width of the interquartile range indicate, respectively, the popularity and the consistency of the themes in each period analyzed.

Figure 5 - Trend Topics.



Note: The Trend Topics chart covers the period from 2018 to 2024.

Source: prepared by the authors.

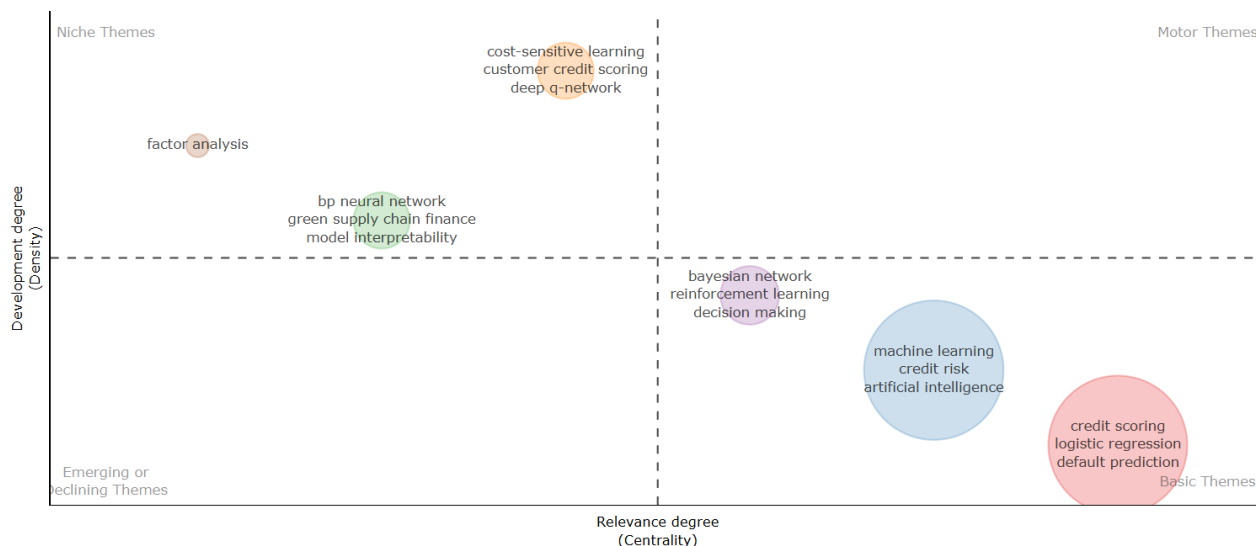
From 2022 onward, there is a noticeable rise in themes related to artificial intelligence and ML, particularly artificial intelligence, ML, deep learning, explainable AI, and XGBoost. This development reflects the global diffusion of AI technologies and their growing application in analytical and decision-making contexts. At the same time, terms associated with credit risk and finance, such as credit risk

model, risk assessment, credit default, and credit rating, remain prevalent, indicating the continued scientific interest in financial risk prediction and management models.

The convergence between AI and finance has become a dominant axis, incorporating ML algorithms into credit evaluation, fraud detection, and predictive modeling. The prominence of explainable AI underscores the concern with transparency and auditability. The term supply chain also emerges, linked to studies on logistical resilience, while artificial neural networks demonstrate the persistence of classical approaches. In summary, the field shows scientific and interdisciplinary maturity, oriented toward the pursuit of models that are more accurate, interpretable, and applicable to complex economic contexts.

Finally, the strategic map (Figure 6) organizes the research themes in a two-dimensional framework that combines centrality, an indicator of relevance and interconnection, with density, which expresses the internal development of each theme. This configuration identifies four quadrants: motor themes, which are highly central and well developed; basic themes, which are central but less developed; niche themes, which are well developed but peripheral; and emerging or declining themes, which are both less central and less developed.

Figure 6 - Strategic map.



Source: prepared by the authors.

In the lower-right quadrant, the cluster "Credit Scoring, Logistic Regression, and Default Prediction" represents the classical and central core, grounded in established statistical techniques. Between the quadrants of basic and motor themes, the cluster "Machine Learning, Credit Risk, and Artificial Intelligence" emerges as the most dynamic axis, characterized by high centrality and density. At the center of the map,

the cluster “Bayesian Network, Reinforcement Learning, and Decision Making” integrates probabilistic inference and adaptive decision policies. In the upper-left area, niche themes such as “BP Neural Network, Green Supply Chain Finance, and Model Interpretability” and “Cost-Sensitive Learning, Customer Credit Scoring, and Deep Q-Network” display high density and lower centrality, reflecting methodological specialization in specific contexts. The cluster “Factor Analysis” functions as a complementary technique for diagnosis and variable reduction, without playing a central role in the field’s development. Table 4 presents the main articles, findings, and directions associated with each cluster.

Table 4 - Analysis of the articles.

Cluster	Article	Main findings	Challenges	Future directions
1. Bayesian Network, Reinforcement Learning e Decision Making	Credit risk scoring with Bayesian network models (LEONG, 2016) Bayesian network oriented transfer learning method for credit scoring model (IWAI; AKIYOSHI; HAMAGAMI, 2021) Using reinforcement learning to optimize the acceptance threshold of a credit scoring model (HERASYMOVYCH; MÄRKA; LUKASON, 2019)	Probabilistic and adaptive models advance credit default prediction and credit decision-making. Bayesian networks handle imbalance, censoring, and scalability effectively, outperforming logistic regression and neural networks in accuracy. Bayesian transfer learning improves cross-domain adaptation, while reinforcement learning dynamically adjusts acceptance thresholds, enhancing profitability.	Challenges remain in adapting and generalizing to dynamic environments. The balance between performance and interpretability is constrained by model complexity. Real-time integration and validation with heterogeneous data are still essential for practical application.	Research suggests hybrid models combining Bayesian networks, reinforcement learning, and adaptive transfer learning, along with explainable methods that ensure traceability and performance. Continuous feedback and online learning emerge as trends for more agile decision-making.
2. BP Neural Network, Green Supply Chain Finance e Model Interpretability	Exploration on the financing risks of enterprise supply chain using back propagation neural network (CAI <i>et al.</i> 2020). Assessing the risk of green building materials certification using the back-propagation neural network (ZHANG; ZHANG; JIANG, 2022) Research on green supply chain finance risk identification based on two-stage deep learning (LIU <i>et al.</i> , 2024)	Neural networks, such as BPNN (Backpropagation Neural Network), are effective in assessing financial and environmental risks in supply chains, handling nonlinear relationships. The LMBP (Levenberg–Marquardt Backpropagation) model outperformed conventional BPNN in accuracy, efficiency, and generalization in green material certification. Two-stage deep learning models improved performance metrics by addressing data imbalance and multisource challenges in green finance.	Persistent challenges include model interpretability, handling asymmetric and scarce data, and achieving practical integration between sustainability and finance. The complexity of interfirm relationships and the lack of standardized green indicators affect reliability. Hyperparameter tuning and scalability in real-world datasets also remain difficult.	It is recommended to enhance transparency and interpretability through explainability techniques and cross-validation. Expanding hybrid and multistage models that combine deep learning and statistical methods may improve stability. Strengthening environmental and financial databases and developing standardized green risk metrics are essential to increase practical applicability.

3. Cost-Sensitive Learning, Learning, Customer Credit Scoring e Deep Q-Network	Cost-sensitive boosted tree for loan evaluation in peer-to-peer lending (XIA; LIU; LIU, 2017). Dynamic classifier ensemble model for customer classification with imbalanced class distribution (XIAO <i>et al.</i>, 2012) Deep reinforcement learning with confusion-matrix-based dynamic reward function for customer credit scoring (WANG <i>et al.</i> 2022)	Advanced modeling approaches have improved credit scoring and risk assessment in imbalanced and dynamic data environments. These methods adjust parameters and rewards according to system behavior, leading to greater profitability, higher accuracy in identifying defaulters, faster convergence, and increased stability compared to traditional approaches.	Handling imbalanced data, dynamic changes, and profitability metrics beyond accuracy remains complex. Models such as DQN require large datasets and may be sensitive to parameter tuning. Ensuring robustness and interpretability across diverse credit contexts continues to be a challenge.	Future research should combine cost-sensitive learning, dynamic ensembles, and reinforcement learning to balance accuracy and adaptability. Adjustable reward functions, realistic economic metrics, and validation across multiple contexts can make scoring systems more robust and efficient.
4. Credit Scoring, Logistic Regression e Default Prediction	An experimental comparison of classification algorithms for imbalanced credit scoring data sets (BROWN; MUES, 2012) A comparative assessment of ensemble learning for credit scoring (WANG <i>et al.</i>, 2011) Statistical and machine learning models in credit scoring: a systematic literature survey (DASTILE, CELIK; POTSANE, 2020)	Ensemble techniques outperform individual classifiers in default prediction, demonstrating greater robustness to class imbalance. Logistic regression remains popular for its simplicity and interpretability. Model ensembles achieve superior performance, while deep learning, despite its potential, remains underexplored.	Class imbalance still reduces the accuracy of some models, and the low explainability of complex techniques hinders financial decision-making. Trade-offs between type I and type II errors affect risk assessment, and the application of deep learning in real-world credit contexts remains limited.	Future research should prioritize methods robust to class imbalance and hybrid frameworks that combine performance and interpretability. The application of explainable deep learning and the evaluation of models based on financial metrics and error costs are promising directions.
5. Factor Analysis	Technology credit scoring model considering both SME characteristics and economic conditions: the Korean case (MOON; SOHN, 2010) Application of deep learning neural network in online supply chain financial credit risk assessment (XU; HE, 2020) Research on enterprise credit risk prediction based on text information (ZHANG, 2021)	Factor analysis reduces financial and non-financial indicators in credit risk prediction models. Applications in online supply chains using deep neural networks demonstrate high accuracy, outperforming SVM and logistic regression. Combining financial and textual variables enhances default prediction. In technology-based SMEs, including business attributes and economic conditions increases predictive power and supports credit management.	Integrating financial, textual, and technological data into coherent models is complex. Deep networks require large datasets and have low interpretability, limiting practical applicability. Generalizing models across sectors and economic contexts, particularly for SMEs and supply chains, remains a challenge.	Future studies should integrate factor analysis with advanced ML, combining multiple data sources to improve accuracy and interpretability. Exploring textual variables, adaptive scorecards, and stress testing under different economic scenarios may yield more robust and reliable models for credit decision support.

6. Machine Learning, Credit Risk e Artificial Intelligence	A multiobjective weighted voting ensemble classifier based on differential evolution algorithm for text sentiment classification (ONAN <i>et al.</i>, 2016) Fairness-aware classifier with prejudice remover regularizer (KAMISHIMA, <i>et al.</i>, 2012) Fairness-aware learning through regularization approach (KAMISHIMA; AKAHO; SAKUMA. 2011)	ML techniques, such as multi-objective weighted ensembles, can improve credit risk prediction, outperforming traditional methods. Moreover, incorporating social fairness concerns and mitigating indirect biases through regularizers in discriminative models, such as logistic regression, enhances fairness without compromising accuracy	Balancing predictive performance and fairness is complex, as indirect biases may persist even when sensitive variables are excluded. Additional challenges include scalability, adaptation to different credit contexts, and integrating fairness metrics with traditional performance indicators.	Future research should develop frameworks that systematically integrate accuracy and fairness, exploring ensembles and adaptive regularization. Incorporating social and legal metrics, combined with validation in real-world, diverse datasets, will enable more transparent and ethical credit risk models.
---	---	--	---	---

Source: prepared by the authors.

Table 4 highlights the consolidation and diversification of modeling approaches in the field of default prediction and credit decision-making. A transition is evident between classical statistical techniques, such as logistic regression, and advanced methods based on ML, reinforcement, and neural networks. The clusters reveal a maturing field: while probabilistic and adaptive models (Cluster 1) and neural networks applied to green supply chains (Cluster 2) enhance generalization and predictive accuracy, concerns with interpretability, transparency, and practical integration are also emerging. Cost-Sensitive and dynamic ensemble approaches (Cluster 3) introduce more realistic economic metrics and adaptability to imbalanced data, and the persistence of logistic regression as a conceptual benchmark (Cluster 4) reflects the search for balance between performance and explainability.

Cluster 4 highlights the central role of credit scoring research in the evolution of default prediction studies. The prominence of authors such as Brown and Mues (2012) and Lessmann *et al.* (2015) within this cluster indicates the continued relevance of discussions concerning the comparative performance of traditional statistical techniques and machine learning approaches. The structure of the cluster suggests that the field has evolved toward increasingly sophisticated predictive models while maintaining strong interest in interpretability, benchmarking practices, and model comparison frameworks.

In parallel, there is growing progress in studies that incorporate multiple data sources and ethical dimensions. The integration of financial, textual, and technological data (Cluster 5) increases the robustness and practical relevance of models for small and medium-sized enterprises and supply chains. Research oriented toward algorithmic fairness (Cluster 6) signals a normative and social shift in the use of ML in credit, emphasizing bias mitigation and decision equity. Cluster 6 reflects the growing incorporation of ethical and governance-related concerns into machine learning applications for credit risk assessment. The concentration of studies addressing fairness, bias mitigation, and responsible

AI indicates that contemporary research is extending beyond predictive performance toward broader concerns involving transparency, accountability, and decision legitimacy. The configuration of this cluster suggests that fairness-oriented approaches have become an important research stream within the field, reflecting the increasing demand for trustworthy and socially responsible analytical systems. Overall, the literature converges toward the development of hybrid, interpretable, and fair machine learning approaches capable of balancing accuracy, adaptability, and ethical responsibility, key challenges for the next generation of credit risk models.

5 CONCLUSION

This study aimed to conduct a bibliometric analysis of the international literature on the use of ML techniques in default prediction, with an emphasis on the automotive sector. A total of 1,026 documents published between January 2010, and September 2025 were examined from the Scopus database to identify thematic clusters, emerging trends, and methodological challenges. The results reveal a rapidly consolidating field, characterized by international expansion, thematic diversification, and methodological maturity.

The main findings can be summarized in four dimensions: (i) the consolidation of ML as the core of default prediction, gradually replacing traditional statistical models with more flexible and self-adaptive approaches; (ii) the broadening of methodological perspectives, with emphasis on neural networks, reinforcement learning, and ensemble models, which enhance accuracy and generalization; (iii) the strengthening of the interpretability and algorithmic transparency agenda, reflecting growing concern with traceability and model trustworthiness; and (iv) the integration of multiple data sources, which improves the robustness of predictions and expands their applicability in diverse contexts. In addition, the lack of systematic investigations focused on the use of ML models for default prediction specifically in the automotive sector highlights a relevant gap in the literature and a promising avenue for future research. Taken together, these findings demonstrate how ML has evolved from a complementary analytical tool into a central paradigm for default prediction, shaping both methodological developments and emerging research agendas in the field.

From a theoretical perspective, the bibliometric landscape outlined here helps to understand the evolution of approaches and production networks that structure the field. From a managerial standpoint, the findings provide strategic insights for the development of more robust, transparent, and explainable credit risk assessment practices, supporting the adoption of data-driven and regulation-aligned credit

policies. Thus, the study contributes to consolidating ML as a technical and epistemological foundation for innovation in credit risk management.

Like any bibliometric study, this research is subject to methodological limitations. The exclusive use of the Scopus database, despite its broad coverage, may have limited the inclusion of technical literature, corporate reports, and publications indexed in other scientific repositories. In addition, variations in terminology across studies required the adoption of thesauri and manual standardization procedures, while the construction of bibliometric networks is inherently sensitive to parameter settings, including thresholds, clustering algorithms, and time windows. Although these factors may influence the composition and structure of the analyzed corpus, the procedures adopted sought to ensure consistency, transparency, and reproducibility. Nevertheless, these limitations do not compromise the overall contribution of the study, which provides a comprehensive overview of the evolution, structure, and emerging directions of research on ML applications for default prediction.

Future research may deepen the qualitative analysis of the most influential clusters, particularly those focused on integrating internal corporate data (such as payment histories and ERP records) with external indicators (such as credit scores and macroeconomic variables). Additional studies may also evaluate the application of hybrid ML models in the automotive sector and investigate how the most effective techniques identified in the literature can be operationalized within real-world credit decision-making processes. By providing a comprehensive understanding of the intellectual development of ML-based default prediction and highlighting opportunities within the automotive sector, this study offers a foundation for future research and managerial innovation at the intersection of finance, artificial intelligence, and credit risk management.

REFERENCES

ALTMAN, E. I. Financial ratios, discriminant analysis and the prediction of corporate bankruptcy. **The Journal of Finance**, v. 23, n. 4, p. 589–609, 1968. DOI: 10.2307/2978933.

ARIZA-GARZÓN, M. J.; ARROYO, J.; CAPARRINI, A.; SEGOVIA-VARGAS, M. J. Explainability of a machine learning granting scoring model in peer-to-peer lending. **IEEE Access**, v. 8, p. 64873–64890, 2020. DOI: 10.1109/ACCESS.2020.2984918.

ASSOCIAÇÃO NACIONAL DOS FABRICANTES DE VEÍCULOS AUTOMOTORES. **Coletiva de Imprensa – Setembro 2025: Desempenho da Indústria Automotiva Brasileira**. Anfavea, 2025.

BAHNSEN, A. C.; AOUADA, D.; OTTERSTEN, B. Example-dependent cost-sensitive logistic regression for credit scoring. In: **INTERNATIONAL CONFERENCE ON MACHINE LEARNING AND APPLICATIONS (ICMLA)**, 2014. Anais... p. 263–269. DOI: 10.1109/ICMLA.2014.48.

BELHADI, A.; KAMBLE, S. S.; VENKATESH, V.; BENKHATI, I.; TOURIKI, F. E. An ensemble machine learning approach for forecasting credit risk of agricultural SMEs' investments in agriculture 4.0 through supply chain finance. **Annals of Operations Research**, v. 345, n. 2, p. 779–807, 2025. DOI: 10.1007/s10479-021-04366-9.

BEQUÉ, A.; LESSMANN, S. Extreme learning machines for credit scoring: an empirical evaluation. **Expert Systems with Applications**, v. 86, p. 42–53, 2017. DOI: 10.1016/j.eswa.2017.05.050.

BLANCO, A.; PINO-MEJÍAS, R.; LARA-RUBIO, J.; RAYO, S. Credit scoring models for the microfinance industry using neural networks: evidence from Peru. **Expert Systems with Applications**, v. 40, n. 1, p. 356–364, 2013. DOI: 10.1016/j.eswa.2012.07.051.

BROWN, I.; MUES, C. An experimental comparison of classification algorithms for imbalanced credit scoring data sets. **Expert Systems with Applications**, v. 39, n. 3, p. 3446–3453, 2012. DOI: 10.1016/j.eswa.2011.09.033.

BÜCKER, M.; SZEPANNEK, G.; GOSIEWSKA, A.; BIECEK, P. Transparency, auditability and explainability of machine learning models in credit scoring. **Journal of the Operational Research Society**, v. 73, n. 1, p. 70–90, 2022. DOI: 10.1080/01605682.2020.1861410.

CAI, X.; QIAN, Y.; BAI, Q.; LIU, W. Exploration on the financing risks of enterprise supply chain using back propagation neural network. **Journal of Computational and Applied Mathematics**, v. 367, 112457, 2020. DOI: 10.1016/j.cam.2019.112457.

CALLON, M.; COURTIAL, J. P.; TURNER, W. A.; BAUIN, S. From translations to problematic networks: an introduction to co-word analysis. **Social Science Information**, v. 22, n. 2, p. 191–235, 1983. DOI: 10.1177/053901883022002003.

CHANG, Y.-C.; CHANG, K.-H.; WU, G.-J. Application of extreme gradient boosting trees in the construction of credit risk assessment models for financial institutions. **Applied Soft Computing**, v. 73, p. 914–920, 2018. DOI: 10.1016/j.asoc.2018.09.029.

CHANG, Y.-C.; CHANG, K.-H.; CHU, H.-H.; TONG, L.-I. Decision tree-based short-term default credit risk assessment models. **Communications in Statistics – Theory and Methods**, v. 45, n. 21, p. 6406–6420, 2016. DOI: 10.1080/03610926.2014.991060.

DASTILE, X.; CELIK, T.; POTSANE, M. Statistical and machine learning models in credit scoring: a systematic literature survey. **Applied Soft Computing**, v. 91, 106263, 2020. DOI: 10.1016/j.asoc.2020.106263.

DENZIN, N. K. **The research act: a theoretical introduction to sociological methods**. 2. ed. New York: McGraw-Hill, 1978.

DUMITRESCU, E.; HUÉ, S.; HURLIN, C.; TOKPAVI, S. Machine learning for credit scoring: improving logistic regression with non-linear decision-tree effects. **European Journal of Operational Research**, v. 302, n. 1, p. 273–289, 2022. DOI: 10.1016/j.ejor.2021.06.053.

HERASYMOVYCH, M.; MÄRKA, K.; LUKASON, O. Using reinforcement learning to optimize the acceptance threshold of a credit scoring model. **Applied Soft Computing**, v. 84, 105697, 2019. DOI: 10.1016/j.asoc.2019.105697.

IWAI, K.; AKIYOSHI, M.; HAMAGAMI, T. Bayesian network oriented transfer learning method for credit scoring model. **IEEJ Transactions on Electrical and Electronic Engineering**, v. 16, n. 9, p. 1195–1202, 2021. DOI: 10.1002/tee.23417.

KAMISHIMA, T.; AKAHO, S.; SAKUMA, J. Fairness-aware learning through regularization approach. In: **IEEE 11TH INTERNATIONAL CONFERENCE ON DATA MINING WORKSHOPS**, 2011. Anais... p. 643–650. DOI: 10.1109/ICDMW.2011.83.

KAMISHIMA, T.; AKAHO, S.; ASOH, H.; SAKUMA, J. Fairness-aware classifier with prejudice remover regularizer. **Lecture Notes in Artificial Intelligence**, v. 7524, p. 35–50, 2012. DOI: 10.1007/978-3-642-33486-3_3.

LEONG, C. K. Credit risk scoring with Bayesian network models. **Computational Economics**, v. 47, n. 3, p. 423–446, 2016. DOI: 10.1007/s10614-015-9505-8.

LESSMANN, S.; BAESENS, B.; SEOW, H.-V.; THOMAS, L. C. Benchmarking state-of-the-art classification algorithms for credit scoring: an update of research. **European Journal of Operational Research**, v. 247, n. 1, p. 124–136, 2015.

LIN, S.; SONG, D.; CAO, B.; GU, X.; LI, J. Credit risk assessment of automobile loans using machine learning-based SHapley Additive exPlanations approach. *Engineering Applications of Artificial Intelligence*, v. 147, 110236, 2025. <https://doi.org/10.1016/j.engappai.2025.110236>.

LIU, Y.; LI, S.; YU, C.; LV, M. Research on green supply chain finance risk identification based on two-stage deep learning. **Operations Research Perspectives**, v. 13, 100311, 2024. DOI: 10.1016/j.orp.2024.100311.

MOON, T. H.; SOHN, S. Y. Technology credit scoring model considering both SME characteristics and economic conditions: the Korean case. **Journal of the Operational Research Society**, v. 61, n. 4, p. 666–675, 2010. DOI: 10.1057/jors.2009.7.

MORESI, E. A. D.; PINHO, I. Proposta de abordagem para refinamento de pesquisa bibliográfica. **New Trends in Qualitative Research**, v. 9, p. 11–20, 2021. DOI: 10.36367/ntqr.9.2021.11-20.

ONAN, A.; KORUKOĞLU, S.; BULUT, H. A multiobjective weighted voting ensemble classifier based on differential evolution algorithm for text sentiment classification. **Expert Systems with Applications**, v. 62, p. 1–16, 2016. DOI: 10.1016/j.eswa.2016.06.005.

REUTERS. Brazil interest rates must stay on hold until data converge, says CenBank director. **Reuters**, 2025. Disponível em: <https://www.reuters.com/world/americas/brazil-interest-rates-must-stay-hold-until-data-converge-says-cenbank-director-2025-10-16/>.

SAUNDERS, A.; ALLEN, L. **Credit risk management in and out of the financial crisis: new approaches to value at risk and other paradigms**. 4. ed. Hoboken: Wiley, 2021.

SILVA, J. P.; NEVES, M. F. **Gestão de riscos corporativos: teoria e prática**. São Paulo: Atlas, 2020.

THOMAS, L. C.; CROOK, J. N.; EDELMAN, D. B. **Credit scoring and its applications**. 2. ed. Philadelphia: Society for Industrial and Applied Mathematics, 2017.

WANG, G.; HAO, J.; MA, J.; JIANG, H. A comparative assessment of ensemble learning for credit scoring. **Expert Systems with Applications**, v. 38, n. 1, p. 223–230, 2011. DOI: 10.1016/j.eswa.2010.06.048.

WANG, Y.; JIA, Y.; TIAN, Y.; XIAO, J. Deep reinforcement learning with the confusion-matrix-based dynamic reward function for customer credit scoring. **Expert Systems with Applications**, v. 200, 117013, 2022. DOI: 10.1016/j.eswa.2022.117013.

WU, Z.; WANG, S.; LI, S. Advances in credit risk modeling using ensemble and deep learning methods. **European Journal of Operational Research**, v. 326, n. 2, p. 357–373, 2025. DOI: 10.1016/j.ejor.2025.04.020.

XIA, Y.; LIU, C.; LIU, N. Cost-sensitive boosted tree for loan evaluation in peer-to-peer lending. **Electronic Commerce Research and Applications**, v. 24, p. 30–49, 2017. DOI: 10.1016/j.elerap.2017.06.004.

XIAO, J.; LI, S.; TIAN, Y.; HUANG, J.; JIANG, X.; WANG, S. Example dependent cost-sensitive learning based selective deep ensemble model for customer credit scoring. **Scientific Reports**, v. 15, n. 1, 2025. DOI: 10.1038/s41598-025-89880-7.

XIAO, J.; XIE, L.; HE, C.; JIANG, X. Dynamic classifier ensemble model for customer classification with imbalanced class distribution. **Expert Systems with Applications**, v. 39, n. 4, p. 3668–3675, 2012. DOI: 10.1016/j.eswa.2011.09.059.

XU, R.; HE, M. Application of deep learning neural network in online supply chain financial credit risk assessment. In: **INTERNATIONAL CONFERENCE ON COMPUTER INFORMATION AND BIG DATA APPLICATIONS (CIBDA)**, 2020. Anais... IEEE, 2020. DOI: 10.1109/CIBDA50819.2020.00058.

ZANCAN, C.; PASSADOR, J. L.; PASSADOR, C. S. Modelos de inteligência artificial na gestão de consórcios intermunicipais brasileiros. *Revista Gestão e Desenvolvimento*, v. 20, n. 2, p. 81–116, 2023. DOI: 10.25112/rgd.v20i2.3424.

ZHANG, C.; ZHANG, J.; JIANG, P. Assessing the risk of green building materials certification using the back-propagation neural network. **Environment, Development and Sustainability**, v. 24, p. 6925–6952, 2022. DOI: 10.1007/s10668-021-01734-0.

ZHANG, H.; ZHANG, H.; ZHANG, M. Research on enterprise credit risk prediction based on text information. **Journal of Risk Analysis and Crisis Response**, v. 11, n. 4, p. 176–188, 2021. DOI: 10.54560/jracr.v11i4.311.

ZUPIC, I.; ČATER, T. Bibliometric methods in management and organization. **Organizational Research Methods**, v. 18, n. 3, p. 429–472, 2015. DOI: 10.1177/1094428114562629.