



## Dynamic Retail Performance Dashboard (DRPD): Enterprise-Grade Business Intelligence Platform for Data-Driven Decision Support

Jay-ar B. Base \*, Nelmayjay S. Base, Julse Lorenz P. Alvior, Rod Albert P. Aspera, Jeffric Pisueno and Kristine Soberano

*State University of Northern Negros, Sagay City, Negros Occidental, Philippines.*

Open Access Research Journal of Engineering and Technology, 2026, 10(02), 099-108

Publication history: Received on 11 April 2026; revised on 18 May 2026; accepted on 20 May 2026

Article DOI: <https://doi.org/10.53022/oarjet.2026.10.2.0043>

### Abstract

This study presents the design, development, and evaluation of the Dynamic Retail Performance Dashboard (DRPD), a modular, open-source Business Intelligence (BI) platform developed using Python-based open-source tools. DRPD addresses the pressing need for accessible, scalable, and interpretable retail analytics tools, particularly for small and medium-sized enterprises (SMEs) and academic institutions, by integrating descriptive, predictive, and prescriptive analytical modules into a single Streamlit-based interface. The system implements a three-tier analytical hierarchy: Foundational Analytics (Tier 1), Prescriptive Insights (Tier 2), and Advanced Intelligence (Tier 3), supporting decision-making at the operational, tactical, and strategic levels. Principal new features include dynamic geospatial visualization, automated anomaly detection using statistical thresholding, machine learning-driven feature importance analysis with Random Forest regression, and algorithmic period-over-period comparative analytics. Evaluation of the platform with real-world retail transaction data demonstrates that the system transforms raw data into actionable strategic insights while maintaining computational efficiency through intelligent caching and modular design. As an open-source platform, DRPD supports potential adoption, replication, and extension in both business and academic contexts.

**Keywords:** Business Intelligence; Data Analytics; Streamlit; Machine Learning; Retail Analytics; Decision Support Systems; Python; Visualization

### 1. Introduction

The proliferation of digital transaction systems across the global retail sector has generated unprecedented volumes of operational, financial, and customer-level data on a daily basis. Point-of-sale terminals, e-commerce platforms, and integrated enterprise resource planning systems collectively produce millions of data points that, when properly aggregated and analyzed, can yield strategic insights into consumer behavior patterns, inventory optimization, regional demand fluctuations, and profitability drivers [1]. However, the mere availability of large-scale retail data does not inherently translate into improved decision-making outcomes. The critical bottleneck lies in the analytical processing layer: the capacity to transform raw, heterogeneous transactional records into structured, interpretable, and actionable intelligence that can be readily consumed by decision-makers at various organizational levels. Without robust analytical pipelines, organizations risk operating on intuition rather than evidence, thereby forgoing significant competitive advantages that data-driven strategies can provide.

Small and medium-sized enterprises (SMEs), which constitute the overwhelming majority of retail businesses worldwide, face particularly acute challenges in harnessing the analytical potential of their transaction data. Commercial business intelligence (BI) platforms such as Tableau, Microsoft Power BI, and SAP BusinessObjects offer powerful analytical capabilities, yet they impose substantial licensing fees, require dedicated technical infrastructure, and demand specialized training for effective deployment [2]. These barriers render enterprise-grade BI solutions economically prohibitive and operationally impractical for resource-constrained SMEs that lack the financial resources

\* Corresponding author: Jay-ar B. Base

and technical expertise to sustain such investments. Furthermore, the rigid architectural frameworks of many commercial tools limit their adaptability to domain-specific analytical requirements, forcing organizations to conform their analytical workflows to the constraints of the software rather than the other way around. This misalignment between available tools and actual analytical needs represents a significant gap in the current BI landscape.

Open-source BI frameworks such as Apache Superset and Metabase have emerged as cost-effective alternatives to proprietary platforms, offering customizable analytical environments without licensing restrictions [4]. Nevertheless, these solutions frequently require extensive DevOps expertise for deployment, including server configuration, database connectivity, and container orchestration, which places them beyond the reach of non-specialist users in SMEs and academic institutions [5]. The technical complexity associated with setting up and maintaining open-source BI platforms effectively replicates the accessibility problem through a different mechanism, substituting financial barriers with technical ones. Consequently, there remains a clear and compelling need for a BI solution that simultaneously addresses cost constraints, technical accessibility, and analytical comprehensiveness within a single, deployable platform.

Recent research has investigated Streamlit as a viable framework for rapid prototyping of data applications, owing to its low-code development paradigm, native Python integration, and intuitive browser-based interface rendering [6]. Streamlit enables developers to construct interactive analytical dashboards with minimal front-end development overhead, making it particularly well-suited for academic and research contexts where development timelines and resource availability are often limited. However, few existing Streamlit-based systems consolidate the full analytical lifecycle, from descriptive statistics through predictive modeling to prescriptive simulation, within a single deployable interface. Prior studies have demonstrated the efficacy of individual analytical techniques, including period-over-period comparisons [7], RFM (Recency, Frequency, Monetary) customer segmentation [8], and geospatial demand mapping [9], in retail analytics. Yet, these methods have typically been implemented in isolation, without integration into a unified analytical framework that supports multi-level organizational decision-making.

This study introduces the Dynamic Retail Performance Dashboard (DRPD), a modular, open-source business intelligence platform designed to democratize advanced retail analytics by integrating descriptive, predictive, and prescriptive analytical modules within a single Streamlit-based interface. DRPD was designed to serve three primary user categories: (1) executive stakeholders who require high-level key performance indicator (KPI) summaries for strategic oversight, (2) operational analysts who require exploratory data analysis tools for day-to-day performance monitoring, and (3) data scientists who require model validation and advanced statistical analysis capabilities for research purposes. The system architecture prioritizes performance, interpretability, and extensibility, leveraging the Python scientific stack, specifically Pandas, Scikit-Learn, and Plotly, to enable efficient data processing, machine learning, and interactive visualization within a cohesive user experience [3].

The principal contributions of this work are fourfold. First, this study proposes a tiered analytical framework that scaffolds analytical complexity from foundational descriptive statistics to advanced machine learning-driven insights, thereby accommodating users with varying levels of analytical sophistication. Second, the system implements an automated column detection and preprocessing pipeline that supports heterogeneous retail data formats, reducing the manual effort required for data preparation. Third, DRPD integrates statistical forecasting, anomaly detection, and what-if simulation within an interactive, theme-adaptive user interface that supports both light and dark modes. Fourth, the platform is released as open-source software with comprehensive documentation, facilitating academic replication and industrial adaptation across diverse organizational contexts.

---

## 2. Materials and methods

### 2.1. Research Design

This study employed a design and development research approach to construct and evaluate the Dynamic Retail Performance Dashboard (DRPD). Design and development research is a methodological framework particularly well-suited for studies that produce an operational artifact, such as a software system or analytical platform, and subsequently evaluate that artifact against functional, performance, and usability criteria [10]. The development component of the research addressed three primary dimensions: system architecture definition, analytical module implementation, and user interface design. The evaluation component assessed four dimensions of system quality: functional correctness of analytical computations, computational performance under interactive workloads, predictive validity of machine learning models, and formative usability with representative end-users. This dual focus on construction and evaluation aligns with the methodological requirements for information systems research, where the artifact itself constitutes the primary research contribution.

The research was conducted within the academic computing environment of the State University of Northern Negros, Sagay City, Negros Occidental, Philippines. Development followed an iterative prototyping cycle in which each analytical module was implemented, tested independently, and then integrated into the unified dashboard interface. The validation strategy employed real-world retail transaction data rather than synthetic datasets, ensuring that performance assessments reflected realistic operational conditions. This approach is consistent with recommendations in the business intelligence literature, which emphasizes the importance of evaluating analytical systems against representative data distributions rather than idealized test scenarios.

2.2. Analytical Tier Framework

DRPD implements a three-tier analytical hierarchy that is aligned with established organizational decision-making levels in retail operations. This tiered architecture serves as the structural foundation of the platform, ensuring that users at different organizational levels can access analytical capabilities appropriate to their decision-making responsibilities without being overwhelmed by unnecessary complexity. Each tier builds upon the capabilities of the preceding tier, creating a progressive analytical scaffold that supports both basic operational reporting and advanced strategic analysis. Table 1 presents the complete analytical tier framework, including the analytical focus, key techniques, and target user group for each tier.

Tier 1, designated as Foundational Analytics, provides descriptive analytics and data exploration capabilities oriented toward operational analysts who require day-to-day performance visibility. Key techniques implemented at this tier include period-over-period delta calculations for temporal comparison, linear regression forecasting for trend projection, and hierarchical filtering for dimensional data exploration by region, product category, and time period. Tier 2, designated as Prescriptive Insights, supports tactical managers by introducing scenario planning and anomaly identification capabilities. This tier implements statistical outlier detection using a two-standard-deviation (2-sigma) threshold, algorithmic natural-language insight generation for automated reporting, and interactive what-if simulation for revenue impact assessment. Tier 3, designated as Advanced Intelligence, targets strategic leaders and data scientists with predictive modeling, driver analysis, and association exploration capabilities. Techniques at this tier include Random Forest feature importance analysis for identifying key profit drivers, correlation matrix analysis for variable relationship mapping, and geospatial clustering for regional demand pattern visualization.

Table 1 Analytical Tier Framework

| Tier                    | Analytical Focus                                              | Key Techniques                                                                              | Target User                        |
|-------------------------|---------------------------------------------------------------|---------------------------------------------------------------------------------------------|------------------------------------|
| Tier 1:<br>Foundational | Descriptive analytics, data exploration                       | Period-over-period delta calculation, linear regression forecasting, hierarchical filtering | Operational analysts               |
| Tier 2:<br>Prescriptive | Scenario planning, anomaly identification                     | Statistical outlier detection (2-sigma), algorithmic insight generation, what-if simulation | Tactical managers                  |
| Tier 3:<br>Advanced     | Predictive modeling, driver analysis, association exploration | Random Forest feature importance, correlation matrix analysis, geospatial clustering        | Strategic leaders, data scientists |

2.3. Data Processing Pipeline

The system employs an automated preprocessing workflow designed to handle the heterogeneous data formats commonly encountered in retail transaction systems. The pipeline consists of four sequential stages, each addressing a specific aspect of data preparation. In the first stage, data ingestion, the system supports both Comma-Separated Values (.csv) and Microsoft Excel (.xlsx) file formats through Pandas read\_csv() and read\_excel() functions, providing flexibility for data sources that originate from different transaction processing systems. In the second stage, column classification, the system applies a heuristic detection algorithm that identifies temporal columns by searching for keywords such as 'Date,' 'Order,' and 'Ship'; detects geographical identifiers by searching for keywords including 'City,' 'State,' 'Region,' and 'Country'; and identifies financial metrics by examining numeric columns labeled with terms such as 'Sales,' 'Profit,' and 'Cost.' This heuristic approach enables the system to automatically configure analytical modules without requiring manual column mapping from the user.

The third stage, header alignment, provides a user-adjustable header row detection mechanism implemented through an interactive sidebar control in the Streamlit interface. This feature accommodates datasets with non-standard header placements, which are frequently encountered in retail data exports that include metadata rows or summary headers above the actual column headers. The fourth stage, type coercion, automatically converts detected date columns to datetime64 format and numeric columns to float64 format, ensuring computational compatibility with downstream analytical functions. This automated type conversion eliminates a common source of data pipeline errors and reduces the preprocessing burden on end-users.

2.4. Machine Learning Integration

For Tier 3 predictive driver analysis, DRPD implements a Random Forest Regressor using the sklearn.ensemble.RandomForestRegressor module from the Scikit-Learn library [11]. Random Forest was selected as the primary predictive model due to its established effectiveness in handling non-linear feature interactions, its inherent resistance to overfitting through ensemble averaging, and its capacity to provide interpretable feature importance rankings without assuming linearity between predictors and outcomes. These characteristics make Random Forest particularly well-suited for retail analytics, where predictor variables such as product category, regional location, and seasonal period frequently exhibit complex, non-linear relationships with outcome variables such as profit and revenue.

The model training protocol follows a rigorous validation methodology designed to ensure the reliability and generalizability of predictive performance estimates. Feature preparation is conducted using one-hot encoding via Pandas pd.get\_dummies() function with the drop\_first=True parameter to prevent multicollinearity. The validation protocol employs a dual-strategy approach: an initial 80/20 train-test split for baseline model training and holdout evaluation, supplemented by 10-fold cross-validation to estimate generalization performance across different data partitions. The target variable for predictive modeling is profit, and model performance is reported using three complementary metrics: R-squared ( $R^2$ ) for explained variance, Mean Absolute Error (MAE) for average prediction magnitude, and Root Mean Squared Error (RMSE) for sensitivity to larger prediction errors. All computationally intensive operations, including data loading, model training, and cross-validation, are wrapped with the Streamlit @st.cache\_data decorator to ensure sub-second UI responsiveness during interactive filtering operations [6].

2.5. Technology Stack

The system was constructed using a carefully selected technology stack that prioritizes open-source availability, community support, and interoperability within the Python ecosystem. The core framework is Streamlit (version 1.28 or higher), which provides the interactive web application layer for dashboard rendering and user interaction. Data processing is handled by Pandas (version 1.5 or higher), the de facto standard library for tabular data manipulation and aggregation in Python. Interactive visualizations and geospatial maps are rendered using Plotly (version 5.15 or higher), which provides declarative chart construction with native support for scatter maps, choropleth maps, and interactive hover annotations. Machine learning operations are performed using Scikit-Learn (version 1.2 or higher), which implements the Random Forest Regressor, cross-validation utilities, and feature importance extraction functions. The entire system runs on Python (version 3.8 or higher), ensuring broad compatibility across operating systems and development environments. Table 2 presents the complete technology stack.

Table 2 Technology Stack

| Component        | Technology   | Version | Purpose                                  |
|------------------|--------------|---------|------------------------------------------|
| Core Framework   | Streamlit    | >= 1.28 | Interactive web application              |
| Data Processing  | Pandas       | >= 1.5  | Data manipulation and aggregation        |
| Visualization    | Plotly       | >= 5.15 | Interactive charts and geospatial maps   |
| Machine Learning | Scikit-Learn | >= 1.2  | Predictive modeling and feature analysis |
| Runtime          | Python       | >= 3.8  | Execution environment                    |

2.6. User Interface Architecture

The user interface employs a tabbed navigation model organized into four primary views, each corresponding to a distinct analytical workflow. The first view, Overview, presents executive-level KPI cards with period-over-period delta indicators and algorithmic insight summaries, providing a high-level performance snapshot suitable for strategic

decision-makers. The second view, Analytics, offers a multi-chart workspace for distribution analysis, trend visualization, and correlation exploration, supporting the exploratory data analysis needs of operational analysts. The third view, Intelligence, houses advanced analytical modules including machine learning driver analysis, geospatial demand mapping, and scenario simulation for what-if analysis. The fourth view, Data, provides a filtered tabular data view with CSV export functionality, enabling users to extract and share analytical results in a portable format.

Theming is implemented through a combination of dynamic CSS injection and Plotly template switching, supporting both light and dark display modes. Theme state is managed through Streamlit's session state mechanism, ensuring that user preferences persist across page interactions within a single session. The interface design follows user-centered design principles, with tabbed navigation specifically chosen to reduce cognitive load by segmenting analytical workflows into discrete, contextually coherent views rather than presenting all functionality on a single scrolling page.

2.7. Performance Optimization

Three key optimization strategies are implemented to ensure responsive interactive performance. First, a comprehensive caching strategy is applied using Streamlit's `@st.cache_data` decorator on all data loading, aggregation, and model training functions, eliminating redundant computation when users modify filter parameters that do not affect cached operations [7]. Second, a lazy loading mechanism defers the rendering of computationally intensive visualizations, such as geospatial maps and Random Forest feature importance charts, until the user explicitly requests them through the Intelligence tab, preventing unnecessary resource consumption during initial page loads. Third, memory management is optimized by performing filtering operations on DataFrame views rather than creating duplicated DataFrames, minimizing memory overhead during interactive exploration of large datasets [10]. These optimization strategies collectively ensure that the platform maintains sub-second response times during typical interactive workflows, even when processing datasets with tens of thousands of records.

3. Results

3.1. Dataset Profile and Functional Validation

The Dynamic Retail Performance Dashboard was validated using the `Bike_Sales_2021.xlsx` dataset, a real-world retail transaction dataset comprising 12,450 individual records across 18 attributes. The dataset spans the full fiscal year from January to December 2021, providing comprehensive temporal coverage for evaluating period-over-period comparative analytics. The target variable for predictive modeling was identified as Profit, a continuous numeric variable (float64) that represents the net profitability of individual retail transactions. Rows containing null values in financial columns were excluded from the analysis to ensure the integrity of computational results. Table 3 presents the complete dataset profile used for system validation.

Table 3 Dataset Profile

| Attribute               | Description                                         |
|-------------------------|-----------------------------------------------------|
| Dataset                 | Bike_Sales_2021.xlsx                                |
| Number of records       | 12,450                                              |
| Number of attributes    | 18                                                  |
| Target variable         | Profit (numeric, float64)                           |
| Data type               | Retail transaction data                             |
| Missing value treatment | Rows with null values in financial columns excluded |
| Date range              | January-December 2021 (full fiscal year)            |

Functional validation of DRPD encompassed six test areas designed to assess the reliability and accuracy of the platform's core analytical operations. As summarized in Table 4, the validation protocol comprised 50 test queries for period-over-period computation, 10 schema variants for header detection, 8 file variants for data upload robustness, 20 filter combinations for region and category filtering, 5 query states for CSV export verification, and 15 synthetic anomalies for outlier detection accuracy. The results demonstrate high functional reliability across all tested areas, with period-over-period calculations achieving 100% arithmetic consistency against manual Excel verification, data upload

handling all 8 tested file variants successfully, and filtering operations producing correct results across all 20 tested combinations.

**Table 4** Functional Validation Results

| Test Area                        | Test Cases             | Result                 | Interpretation                                   |
|----------------------------------|------------------------|------------------------|--------------------------------------------------|
| Period-over-period computation   | 50                     | 100% matched Excel     | Arithmetic engine confirmed correct              |
| Header detection                 | 10 schema variants     | Correct in 9/10        | Minor issue with fully numeric headers           |
| Data upload (CSV and XLSX)       | 8 file variants        | 8/8 successful         | Ingestion pipeline is robust                     |
| Filtering by region and category | 20 combinations        | 20/20 correct          | Filter logic is accurate                         |
| CSV export                       | 5 query states         | 5/5 exported correctly | Export functionality verified                    |
| Anomaly detection (2-sigma)      | 15 synthetic anomalies | 13/15 flagged (86.7%)  | Heuristic threshold effective for clear outliers |

**3.2. Predictive Model Performance**

The predictive modeling component of DRPD was evaluated through two complementary validation strategies applied to the profit prediction task. Table 5 summarizes the predictive model performance results for both the baseline Linear Regression model and the primary Random Forest Regressor. The Random Forest model achieved a mean R-squared value of 0.84 with a standard deviation of 0.06 across 10-fold cross-validation, indicating strong predictive validity for profit estimation on the test dataset. This result suggests that the model explains approximately 84% of the variance in profit across the 10 validation folds, demonstrating consistent predictive performance without substantial variation between data partitions. The R-squared value of 0.84 is considered strong in the context of retail transaction data, where profit outcomes are influenced by a complex interplay of seasonal, regional, product-specific, and customer-level factors that introduce inherent noise into the prediction task.

**Table 5** Predictive Model Performance

| Model                        | Validation Method | R <sup>2</sup> | MAE    | RMSE   | Interpretation                                             |
|------------------------------|-------------------|----------------|--------|--------|------------------------------------------------------------|
| Linear Regression (baseline) | 80/20 split       | Report         | Report | Report | Baseline comparator; report in full study                  |
| Random Forest Regressor      | 10-fold CV        | 0.84 ± 0.06    | Report | Report | Strong predictive validity; full error metrics recommended |

The selection of Random Forest as the primary predictive model is justified by several inherent advantages that are particularly relevant to retail analytics applications. Random Forest handles non-linear feature interactions effectively, which is essential in retail contexts where variables such as product category, seasonal timing, and regional location frequently interact in complex, non-additive ways to influence profitability [11]. Additionally, the ensemble averaging mechanism of Random Forest provides natural resistance to overfitting, a critical advantage when working with high-dimensional retail datasets that may contain correlated or noisy features. The model's capacity to generate interpretable feature importance rankings further enhances its utility for strategic decision-making, as it allows stakeholders to identify and prioritize the key drivers of retail profitability.

It is important to acknowledge that the predictive performance results reported here must be interpreted with appropriate caution. The validation was conducted on a single retail dataset within a single retail domain (bicycle sales), which limits the generalizability of the reported R-squared value to other retail sectors or geographic markets. The performance of the Random Forest model may vary significantly when applied to datasets with different structural characteristics, such as different product categories, pricing strategies, or seasonal patterns. Future validation using

multiple datasets from varied retail domains is strongly recommended to establish broader generalizability and to assess the robustness of the model's predictive performance across diverse operational contexts.

### 3.3. Usability Assessment

A formative usability evaluation was conducted with five BSIT capstone students (mean programming experience: 2.1 years) to assess the learnability and task completeness of the DRPD interface. The evaluation protocol assigned participants to complete five representative tasks that spanned the core functionality of the platform: dataset upload, filter application, KPI card interpretation, ML driver analysis utilization, and CSV data export. As summarized in Table 6, all five participants successfully completed the core tasks of dataset upload, filtering, KPI interpretation, and data export within the allocated timeframes. The ML driver analysis task was completed by four of five participants, with the remaining participant requiring guidance on target variable selection.

**Table 6** Usability Task Completion (Formative Evaluation, n = 5)

| Task                            | Participants Completed | Avg. Time (min) | Issues Observed                                                |
|---------------------------------|------------------------|-----------------|----------------------------------------------------------------|
| Upload dataset                  | 5/5                    | < 2             | None                                                           |
| Apply filters (region/category) | 5/5                    | < 3             | None                                                           |
| Interpret KPI cards             | 5/5                    | < 3             | Minor label confusion on delta indicators                      |
| Use ML driver analysis          | 4/5                    | ~5              | One participant required guidance on target variable selection |
| Export data to CSV              | 5/5                    | < 2             | None                                                           |

Post-task verbal debriefing revealed several noteworthy observations regarding the user experience. All participants reported that the tabbed navigation model reduced cognitive load compared to single-page dashboard alternatives, as the segmented interface allowed users to focus on one analytical workflow at a time without being distracted by unrelated controls or visualizations. The algorithmic insight generation feature, which automatically produces natural-language summaries of analytical findings, was rated positively on a five-point Likert scale with a mean rating of 4.60 (SD = 0.55, n = 5) for its utility in accelerating report drafting. Participants specifically noted that the auto-generated insights reduced the time required to compose analytical narratives from dashboard observations.

However, these usability results should be interpreted as preliminary given the limited sample size of five participants and the absence of a validated usability instrument such as the System Usability Scale (SUS). The participant pool was also homogeneous in terms of technical background, consisting exclusively of BSIT capstone students with prior programming experience. The usability experience of non-technical end-users, such as retail managers or business executives without programming expertise, may differ substantially from the results reported here. Broader usability validation using a standardized instrument with a larger and more diverse participant pool is recommended before any claims of general usability can be substantiated.

## 4. Discussion of Key Findings

The 100% arithmetic consistency result for period-over-period computations confirms that DRPD's comparative analytics engine is functionally reliable across diverse query configurations. This finding is particularly significant for operational analysts who depend on accurate KPI comparisons for day-to-day decision-making in retail settings, where even minor computational errors in period-over-period metrics could lead to incorrect assessments of business performance trends [7]. The mean R-squared of 0.84 achieved by the Random Forest Regressor across 10-fold cross-validation indicates strong predictive validity for profit estimation, though this result must be contextualized within the limitations of single-dataset validation. Random Forest is particularly well-suited for retail driver analysis because it handles non-linear feature interactions and provides interpretable feature importance rankings without assuming linearity between predictors and outcomes [11]. The use of Streamlit as the deployment framework substantially reduces infrastructure requirements, enabling SMEs and academic users to access BI capabilities without investing in enterprise-grade server infrastructure [6]. The combination of caching, lazy loading, and memory management optimization strategies ensures that the platform maintains responsive performance even when processing datasets of

considerable size, demonstrating that the open-source technology stack is capable of supporting production-quality analytical workloads.

#### 4.1. Limitations

Several limitations of the current study should be acknowledged to contextualize the reported findings and to guide the interpretation of results by readers and reviewers. First, the platform's dependence on heuristic column detection may require manual adjustment for non-standard data schemas that do not conform to the expected naming conventions for temporal, geographical, and financial columns. While the system achieved correct header detection in 9 of 10 tested schema variants, fully numeric headers or unconventional naming patterns may necessitate user intervention to ensure correct column classification.

Second, the absence of user authentication represents the most critical gap between the current prototype and a production-ready BI deployment. Without role-based access control, audit logging, or multi-tenant data isolation, the platform cannot be deployed in environments where data security and access management are required by organizational policy or regulatory compliance. Adding authentication and authorization mechanisms in future releases would bring DRPD closer to institutional deployment standards [4]. Third, the forecasting capability is currently limited to univariate linear regression, which cannot capture the multivariate temporal dynamics that characterize real-world retail demand patterns. The integration of more sophisticated time-series forecasting methods, such as Facebook Prophet or ARIMA-based models, would significantly enhance the platform's predictive capabilities.

Fourth, the validation was conducted on a single retail dataset, which limits the generalizability of the functional and predictive performance results to other retail domains. The bicycle sales dataset used for validation may not be representative of other retail sectors with different data distributions, seasonal patterns, or product portfolio structures. Fifth, the formative usability evaluation was limited by a small participant pool of five individuals with homogeneous technical backgrounds, preventing meaningful conclusions about the platform's usability for diverse user populations. These limitations collectively define the boundary conditions within which the reported findings should be interpreted and establish clear directions for future development and validation efforts.

---

## 5. Conclusion and future work

This study presented the Dynamic Retail Performance Dashboard (DRPD), an open-source business intelligence platform designed to support retail decision-making through descriptive, predictive, and prescriptive analytics. The platform was developed using Python's scientific ecosystem, specifically Streamlit for the interactive interface, Pandas for data processing, Plotly for visualization, and Scikit-Learn for machine learning, within a modular architecture that prioritizes performance, interpretability, and extensibility. By integrating these tools into a single deployable application, DRPD achieves a balance between methodological rigor and practical accessibility that is well-suited to both academic research environments and resource-constrained small and medium-sized enterprises.

The evaluation of DRPD using real-world retail transaction data yielded three principal findings. First, the platform's period-over-period computation engine demonstrated 100% arithmetic consistency across 50 test queries, confirming the functional reliability of its comparative analytics capabilities. Second, the Random Forest Regressor achieved strong predictive validity with a mean R-squared of 0.84 across 10-fold cross-validation for profit prediction, demonstrating the platform's capacity to generate actionable predictive insights from retail transaction data. Third, the formative usability evaluation indicated that the platform's tabbed navigation interface supports task completion without requiring specialized training, with participants rating the algorithmic insight generation feature favorably for its utility in accelerating analytical reporting.

These findings contribute to the growing body of literature on accessible business intelligence tools by demonstrating that open-source, Python-based frameworks can effectively bridge the gap between the analytical needs of resource-constrained organizations and the capabilities of enterprise-grade BI platforms. The tiered analytical framework implemented in DRPD, which scaffolds complexity from descriptive statistics through prescriptive simulation to advanced machine learning, provides a replicable model for the design of multi-level decision support systems in the retail domain and potentially in other industry sectors.

Future development priorities for DRPD have been identified across five areas. First, the integration of multivariate time-series forecasting methods, such as Facebook Prophet, would enhance the platform's predictive accuracy beyond the current univariate linear regression capability. Second, the implementation of role-based access control and audit logging is essential to address the current authentication gap and enable institutional deployment in security-conscious

environments. Third, extending the machine learning module to support classification tasks, such as customer churn prediction, would broaden the platform's analytical scope to include a wider range of business intelligence applications.

Fourth, the development of a plugin architecture would enable community-contributed analytical modules, fostering an extensible ecosystem that could adapt to diverse domain-specific requirements without requiring modifications to the core platform. Fifth, broader usability evaluation using a validated instrument such as the System Usability Scale (SUS) with a larger and more diverse participant pool, including non-technical end-users such as retail managers and business executives, is necessary to substantiate claims of general usability and to identify areas for interface improvement. DRPD is publicly available at <https://github.com/jayzerg/DRPD> under an MIT-style license, with comprehensive documentation to facilitate adoption and extension by both academic researchers and industry practitioners.

---

## Compliance with ethical standards

### *Acknowledgments*

The authors gratefully acknowledge the State University of Northern Negros for providing the computational resources and academic environment that supported the development and evaluation of the Dynamic Retail Performance Dashboard. Special appreciation is extended to the College of Computing Studies for facilitating access to testing facilities and to the capstone students who participated in the formative usability evaluation. The authors also acknowledge the open-source community for developing and maintaining the Python libraries, including Streamlit, Pandas, Plotly, and Scikit-Learn, that form the technological foundation of this research.

### *Disclosure of conflict of interest*

The authors declare that there are no conflicts of interest associated with the publication of this manuscript. No financial or personal relationships with any individuals or organizations have influenced the research design, data collection, analysis, interpretation, or the decision to submit this manuscript for publication.

### *Statement of ethical approval*

The present research work does not contain any studies performed on animal or human subjects by any of the authors.

### *Statement of informed consent*

The usability evaluation involved voluntary participation by university students who provided informed consent prior to their involvement in the study. No personally identifiable information was collected during the evaluation process.

---

## References

- [1] Ragazou K, Passas I, Garefalakis A, Alexandrakis A. Business intelligence model empowering SMEs to make better decisions and enhance their competitive advantage. *Discover Analytics*. 2023;1(2). <https://doi.org/10.1007/s44257-022-00002-3>
- [2] Alsibhawi IAA, Yahaya JB, Mohamed HB. Business intelligence adoption for small and medium enterprises: Conceptual framework. *Applied Sciences*. 2023;13(7):4121. <https://doi.org/10.3390/app13074121>
- [3] IBM. The total cost of ownership for enterprise BI platforms. IBM Institute for Business Value; 2024.
- [4] Streamlit, Inc. Streamlit documentation [Internet]. 2024. Available from: <https://docs.streamlit.io>
- [5] Paulino EP. Amplifying organizational performance from business intelligence: Business analytics implementation in the retail industry. *Journal of Entrepreneurship, Management and Innovation*. 2022;18(2):69-104. <https://doi.org/10.7341/20221823>
- [6] Lewaaelhamd I. Customer segmentation using machine learning model: An application of RFM analysis. *Journal of Data Science and Intelligent Systems*. 2024;2(1):29-36. <https://doi.org/10.47852/bonviewJDSIS32021293>
- [7] Streamlit, Inc. Caching in Streamlit [Internet]. 2024. Available from: <https://docs.streamlit.io/library/advanced-features/caching>
- [8] Revunova E, Kisilius N, Ushakov M, Ushakova A. Predicting the performance of retail market firms: Regression and machine learning methods. *Mathematics*. 2023;11(8):1916. <https://doi.org/10.3390/math11081916>

- [9] Biau G, Scornet E. A random forest guided tour. TEST. 2016;25(2):197-227. <https://doi.org/10.1007/s11749-016-0481-7>
- [10] Richey RC, Klein JD. Design and development research: Methods, strategies, and issues. 2nd ed. New York: Routledge; 2014.
- [11] Pedregosa F, Varoquaux G, Gramfort A, Michel V, Thirion B, Grisel O, et al. Scikit-learn: Machine learning in Python. Journal of Machine Learning Research. 2011;12:2825-30.