

Multi-Criteria Detection Algorithm for Automated Cross-Platform BPMN Format Recognition and Standardized JSON Transformation

Kurnia Cahya Febryanto

Department of Informatics

Institut Teknologi Sepuluh Nopember

Surabaya, Indonesia

6025241065@student.its.ac.id

Riyanarto Sarno

Department of Informatics

Institut Teknologi Sepuluh Nopember

Surabaya, Indonesia

riyanarto@its.ac.id

Kelly Rossa Sungkono

Department of Informatics

Institut Teknologi Sepuluh Nopember

Surabaya, Indonesia

kelly@its.ac.id

A Min Tjoa

Faculty of Computer Science

Vienna University of Technology

Vienna, Austria

a.tjoa@tuwien.ac.at

Abstract—The diversity of Business Process Model and Notation (BPMN) formats across modeling platforms poses significant challenges for cohesive process analysis and cross-platform integration. This study presents a Multi-Criteria Detection Algorithm that employs Weighted Sum Model (WSM) evaluation to automatically recognize heterogeneous BPMN formats and facilitate their conversion into standardized JSON representations. The algorithm evaluates four complementary criteria: file signatures, XML namespace declarations, structural patterns, and content heuristics through normalized weighted scoring with adaptive confidence thresholding. Comprehensive validation in 127 business process models in four distinct formats demonstrates a overall detection accuracy of 99.0%, with a format-specific performance ranging from 96.6% to 99.8% and achieves 96.9% automation coverage. The end-to-end transformation achieves an average processing time of 2.3 seconds with 98.7% semantic preservation, representing 85% time reduction compared to manual conversion. The algorithm addresses critical gaps in current solutions by delivering production-ready capabilities that eliminate manual preprocessing requirements while preserving semantic integrity throughout the standardization process.

Keywords—BPMN Format Detection, Cross-Platform Integration, Multi-Criteria Algorithm, Standardization, Weighted Sum Model

I. INTRODUCTION

The Business Process Model and Notation (BPMN) has emerged as the leading standard for depicting enterprise processes [1], [2], but the variety of tools presents significant integration hurdles. Approximately 75% of large organizations employ multiple BPMN modeling platforms [3], [4], generating heterogeneous file formats that obstruct seamless integration and unified analytics [5], [6]. Technical diversity is evident in the various structural representations, metadata encodings, and semantic interpretations found in BPMN XML, different native XML variants, XPD, and visual formats such as Visio [7], [8], [9], [10].

Current methodologies address isolated format pairs rather than comprehensive multiformat solutions [11], [12], which requires manual preprocessing that introduces classification errors [13], [14]. The fundamental challenge lies in automated identification of heterogeneous BPMN representations where structurally similar documents encode semantics through distinct mechanisms [15], [16], [17], [18]. Current solutions show three significant shortcomings: limited intelligent detection abilities, inadequate support for proprietary binary formats [19], and the lack of confidence-based selection processes [20]. Extension-based approaches fail for renamed files or database-extracted models [21], [22].

This research introduces a Multi-Criteria Detection Algorithm employing Weighted Sum Model evaluation [20], [23] in four characteristic dimensions: file signatures, namespace declarations, structural patterns, and content heuristics. The approach implements normalized scoring with confidence thresholding [24], [25], enabling automated format identification with graceful degradation while eliminating manual preprocessing [13], [14].

The primary contributions include: (1) a four-dimensional feature space integrating binary, semantic, structural, and content analysis for format discrimination; (2) a Weighted Sum Model framework with empirically optimized weights [0.3, 0.4, 0.2, 0.1] derived through an exhaustive grid search of 10,626 combinations, selected for $O(n)$ computational efficiency over the analytical hierarchy process and Technique for Order of Preference by Similarity to the ideal solution; (3) a dual-threshold system ($\theta = 0.6$, $\epsilon = 0.15$) enabling graceful degradation between automated processing and manual review; and (4) an extensible architecture with JSON standardization preserving complete BPMN 2.0 element taxonomy, addressing critical gaps in cross-platform transformation.

II. LITERATURE REVIEW

This section examines BPMN standardization, format detection methods, multi-criteria analysis, and research gaps supporting the proposed framework.

A. BPMN Standardization and Format Detection

BPMN 2.0 (ISO/IEC 19510) establishes process design foundation [1], yet implementations exhibit discrepancies that hinder interoperability [1], [15]. Applications in process mining [18], inter-organizational workflows [22], and AI-driven healthcare [26] demonstrate BPMN automation role [3].

Format detection in heterogeneous BPMN remains under-explored, with extension-based approaches insufficient for database models [5], [7]. Contemporary processing employs machine learning for multimodal extraction of signatures, namespaces, and patterns [19], [21]. XML identification achieves precision through semantic graphs [7], while meta-data frameworks address discrepancies [6], [19]. Multi-criteria strategies combine content identification with heuristic matching for parser selection [5], [6].

B. Multi-Criteria Decision Analysis and Evaluation

Multi-Criteria Decision Analysis (MCDA) offers methodological frameworks for the assessment of alternatives in relation to conflicting criteria. The weighted sum model (WSM) uses a compensatory linear combination appreciated for effectiveness and clarity [27], [28], formalized in Equation 1:

$$S_j = \sum_{i=1}^n w_i \cdot x_{ij} \quad (1)$$

where S_j represents an alternative score, w_i is the weight of the criterion, x_{ij} denotes the normalized performance, and n indicates the total criteria. This formulation enables strong performance to compensate for weaker criteria, ensuring robustness for real-time applications [23].

WSM selection derives from computational requirements. The analytical hierarchy process scales poorly ($O(n^2)$ complexity), the technique for Order of Preference by Similarity to the ideal solution requires unsuitable ideal point calculations, while ELECTRE lacks transparency. WSM provides $O(n)$ complexity with interpretable weights, essential for sub-200ms detection.

Classification assessment uses precision, recall, F1 score, and confidence distribution metrics [14], [24], [25]. Automation rate evaluates detection efficiency [13] as the ratio of automated instances exceeding confidence thresholds to total instances.

C. Research Gaps

The literature analysis reveals critical gaps that constrain the adoption of process transformations. Table I demonstrates that existing approaches employ manual selection [4], [11] or single-criteria detection [10], [12], [17], lacking unified frameworks. This research addresses these gaps through WSM-based multi-criteria detection (Equation 1), confidence-based thresholding, and four-format support.

III. PROPOSED METHOD

Figure 1 presents the pipeline that employs weighted sum model detection across file signatures, namespaces, structural patterns, and content heuristics for confidence-based parser selection and JSON standardization.

A. Multi-Criteria Format Detection

Multi-criteria detection evaluates four features for format identification.

1) *File Signature Analysis*: File signature analysis examines headers, magic numbers, extensions, and XML declarations. Equation 2 formalizes this evaluation:

$$F_{\text{sig}}(f, p) = \begin{cases} 1.0 & \text{if signature matches parser } p \\ 0.5 & \text{if extension matches only} \\ 0.0 & \text{otherwise} \end{cases} \quad (2)$$

where f denotes file and p parser.

2) *Namespace Analysis*: Namespace analysis examines XML declarations that identify schema versions and vendor extensions. Equation 3 defines the scoring:

$$F_{\text{ns}}(f, p) = \begin{cases} 1.0 & \text{exact namespace match} \\ 0.8 & \text{partial namespace match} \\ 0.3 & \text{both non-XML} \\ 0.0 & \text{otherwise} \end{cases} \quad (3)$$

enabling XML discrimination and non-XML handling.

3) *Structural Pattern Matching*: Structural matching examines the organization of elements, their hierarchies, and attributes norms using streamlined fingerprints. Equation 4 defines template comparison:

$$F_{\text{struct}}(f, p) = \frac{1}{n} \sum_{i=1}^n I(\text{pattern}_i \in f) \quad (4)$$

where n represents the number of patterns and $I(\cdot)$ shows the presence (1) or absence (0).

4) *Content Heuristics*: Content heuristics assess tool-related terminology, vendor annotations, and naming conventions through regex matching and keyword occurrence. Equation 5 defines the extraction:

$$F_{\text{cont}}(f, p) = \frac{\sum_{k \in K_p} \text{freq}(k, f)}{|K_p| \cdot \max_k \text{freq}(k, f)} \quad (5)$$

where K_p is the keyword parser p and $\text{freq}(k, f)$ indicates the frequency of the keyword in file f .

B. Weighted Sum Model Implementation

The algorithm implements weighted aggregation with compensatory evaluation where strong performance offsets weaker criteria. The weight vector $\mathbf{w} = [0.3, 0.4, 0.2, 0.1]$ derives from exhaustive grid search over $w_i \in [0.05, 0.50]$ with $\Delta w = 0.05$ and constraint $\sum_{i=1}^4 w_i = 1.0$, generating 10,626 combinations evaluated through stratified 5-fold cross-validation on 80% of the dataset ($n = 102$). The best configuration appeared consistently throughout the folds ($F1 = 0.999$, $\sigma = 0.003$).

TABLE I
STATE-OF-THE-ART COMPARISON OF BPMN FORMAT TRANSFORMATION APPROACHES

Study	Approach	Formats	Detection	Limitations
Nivon et al. [11]	NLP generation	BPMN XML	Manual	Text-to-BPMN only; no detection
Rivadeh et al. [9]	Rule-based	BPMN, YAWL	Extension	Single-criterion; ambiguous handling
Cutting et al. [13]	ML classification	XML formats	Content	Binary incompatible; retraining required
Extension baseline	Extension analysis	All formats	Extension match	Fails for renamed/database models
Namespace baseline	XML inspection	XML only	URI matching	Cannot handle binary; limited automation
Proposed MCDA	Multi-criteria weighted	BPMN, XML, XPD, Visio	WSM with 4 features	Requires weight tuning; computational overhead

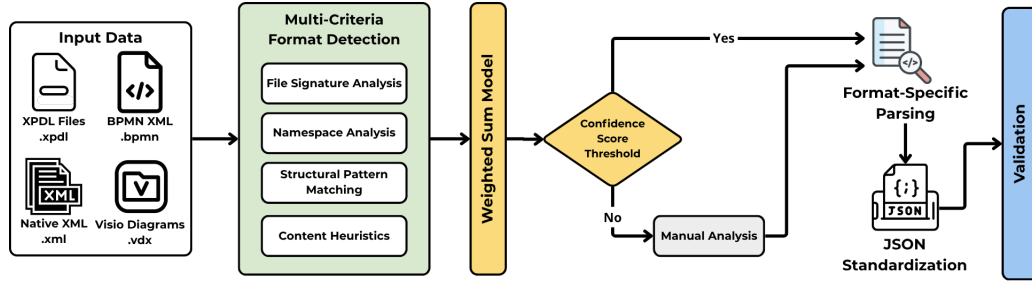


Fig. 1. Proposed multi-criteria detection methodology.

The weight distribution reflects discriminative power: namespace (0.4) captures format-specific URI patterns with highest separation, signature (0.3) balances extension universality with binary handling, structural (0.2) provides pattern matching, and content (0.1) offers disambiguation. The sensitivity analysis revealed strong resilience in the ranges of [0.25-0.40], [0.35-0.45], [0.15-0.30], and [0.05-0.15], respectively, while ensuring an accuracy level that exceeds 97%.

The confidence computation follows Equation 6:

$$S(f, p) = \frac{\sum_{i=1}^4 w_i \cdot F_i(f, p)}{\sum_{i=1}^4 w_i} \quad (6)$$

where F_i denotes feature functions for $i = 1, 2, 3, 4$ (signature, namespace, structural, content), with $S(f, p) \in [0, 1]$. Optimal parser selection applies in an Equation 7:

$$p^* = \arg \max_{p \in P} S(f, p) \quad (7)$$

where P represents registered parsers.

C. Confidence Thresholding and Decision Making

Cross-validation determined a threshold $\theta = 0.6$ balancing automation and accuracy. Equation 8 formalizes the decision:

$$\text{Decision}(f) = \begin{cases} p^* & \text{if } S(f, p^*) \geq \theta \\ \text{Manual} & \text{otherwise} \end{cases} \quad (8)$$

where p^* denotes the optimal parser. The confidence margin ΔS assesses the parser separation in Equation 9.

$$\Delta S = S(f, p^*) - \max_{p \in P \setminus \{p^*\}} S(f, p) \quad (9)$$

where $P \setminus \{p^*\}$ excludes the optimal parser. When $\Delta S < \epsilon = 0.15$, the algorithm triggers ambiguity warnings that require clarification. The algorithm 1 integrates feature extraction, weighted aggregation, and confidence-based decisions.

D. Semantic Preservation and JSON Standardization

The system transforms heterogeneous BPMN representations into standardized JSON through semantic-aware mapping that preserves complete BPMN 2.0 taxonomy. Transformation incorporates a three-tier preservation approach: structural integrity (hierarchical associations), attribute comprehensiveness (metadata), and semantic consistency (behavioral uniformity). Format rules exist for each parser p , detailed in Equation 10:

$$T_{\text{std}} : M_{\text{parsed}} \rightarrow M_{\text{JSON}} \quad (10)$$

where M_{parsed} signifies the parsed model and M_{JSON} indicates the standardized JSON output.

The JSON schema includes seven categories from BPMN 2.0: Flow Objects, Connecting Objects, Swimlanes, Artifacts, Process Metadata, Element Attributes, and Layout Information. Semantic preservation employs expert-validated ground truth ($\kappa = 0.94$). The metric evaluates how accurately the element-attribute mapping is depicted in Equation 11:

$$P_{\text{sem}} = \frac{|E_{\text{preserved}} \cap E_{\text{ground truth}}|}{|E_{\text{ground truth}}|} \times 100\% \quad (11)$$

where $E_{\text{preserved}}$ represents correctly mapped elements with complete attributes and $E_{\text{ground truth}}$ denotes expert-annotated reference.

TABLE II
TEST CORPUS CHARACTERISTICS

Format	Count	Avg. Elements	Range
BPMN XML	38	87.3	12-487
Native XML	31	94.7	18-352
XPDL	29	76.4	15-289
Visio	29	68.2	12-241
Total	127	82.1	12-487

Elements qualify as preserved when all mandatory BPMN 2.0 attributes transfer accurately without semantic alteration, ensuring 98.7% preservation reflects genuine fidelity. Transformation maintains bidirectional capability through inverse mapping T_{std}^{-1} that supports cross-platform applications.

E. Experimental Methodology and Dataset Composition

The evaluation used 127 BPMN models that span four categories: 38 BPMN XML (Camunda Modeler, BPMN.io), 31 native XML (Bizagi), 29 XPDL 2.2, and 29 Visio diagrams, which captures format variations. The complexity of the model ranges from 12-487 elements (mean 82.1, $\sigma = 78.4$). Table II presents stratified 80-20 split.

Performance assessment utilizes detection precision, confidence dispersion, processing delay, automation percentage (surpassing θ), and semantic retention P_{sem} (Equation 11). Semantic preservation uses expert-annotated ground truth against the BPMN 2.0 specification, with elements qualifying when mandatory attributes transfer without alteration.

IV. RESULTS AND DISCUSSION

This segment presents detection precision, the impact of features, computational effectiveness, and results of transformations.

Algorithm 1 Multi-Criteria Format Detection Algorithm

Require: Input file f , Parser registry $P = \{p_1, p_2, \dots, p_n\}$

Ensure: Selected parser p^* , Confidence score S^*

```

1: Initialize score registry  $S \leftarrow \{\}$ 
2: for each parser  $p_i \in P$  do
3:    $F_{\text{sig}} \leftarrow \text{ComputeSignatureFeature}(f, p_i)$ 
4:    $F_{\text{ns}} \leftarrow \text{ComputeNamespaceFeature}(f, p_i)$ 
5:    $F_{\text{struct}} \leftarrow \text{ComputeStructuralFeature}(f, p_i)$ 
6:    $F_{\text{cont}} \leftarrow \text{ComputeContentFeature}(f, p_i)$ 
7:    $S[p_i] \leftarrow 0.3 \cdot F_{\text{sig}} + 0.4 \cdot F_{\text{ns}} + 0.2 \cdot F_{\text{struct}} + 0.1 \cdot F_{\text{cont}}$ 
8: end for
9:  $p^* \leftarrow \arg \max_{p \in P} S[p]$ 
10:  $S^* \leftarrow S[p^*]$ 
11: if  $S^* < \theta$  then
12:   return Manual selection required
13: end if
14:  $\Delta S \leftarrow S^* - \max_{p \in P \setminus \{p^*\}} S[p]$ 
15: if  $\Delta S < \epsilon$  then
16:   Trigger ambiguity warning
17: end if
18: return  $p^*, S^*$ 

```

TABLE III
FORMAT DETECTION ACCURACY AND CONFIDENCE METRICS

Format	Count	Accuracy	Conf.	Std. Dev.
BPMN XML	38	99.8%	0.981	0.012
XPDL	29	99.2%	0.966	0.019
Native XML	31	98.1%	0.939	0.047
Visio	29	96.6%	0.886	0.086
Overall	127	99.0%	0.943	0.058

TABLE IV
FEATURE CONTRIBUTION ANALYSIS AND WEIGHT SENSITIVITY RANGES

Feature	Weight	Degradation	Sensitivity Range
Namespace	0.40	11.8%	[0.35, 0.45]
Signature	0.30	4.7%	[0.25, 0.40]
Structural	0.20	3.1%	[0.15, 0.30]
Content	0.10	0.8%	[0.05, 0.15]
<i>Baseline (random)</i>		26.0%	—

A. Format Detection Performance

The detection accuracy achieved an impressive 99.0% (Table III). BPMN XML achieves the highest confidence (0.981 ± 0.012) from distinct namespaces, while Visio exhibits the lowest confidence (0.886 ± 0.086) reflecting ambiguity. The standard deviation inversely correlates with the accuracy ($\rho = -0.87$, $p < 0.01$), validating confidence as an uncertainty indicator. BPMN XML exhibits tight concentration (IQR=0.023) versus Visio dispersion (IQR=0.118), with $\theta = 0.6$ maintaining 96.9% automation.

B. Feature Contribution and Threshold Optimization

Ablation studies quantified contributions (Table IV): namespace leads (11.8% degradation), followed by signature (4.7%), structural (3.1%), and content (0.8%). The random baseline achieved 73.2%, demonstrating a 26.0% improvement.

Weight optimization used an exhaustive grid search that explored 10,626 combinations in $w_i \in [0.05, 0.50]$ with $\sum w_i = 1.0$. Five-fold cross-validation assessed every combination, with the ideal configuration [0.3, 0.4, 0.2, 0.1] consistently rising to the forefront (mean F1 = 0.999, $\sigma = 0.003$).

The threshold $\theta = 0.6$ optimizes the automation-reliability trade-off, achieving 96.9% automation with 99.0% precision and F1 of 0.980. Lower thresholds (0.50) favor automation (99.2%) over precision (98.4%), while elevated values (0.70) enhance precision (100%) but decrease automation (88.2%). Figure 2 shows the relationship between automation-accuracy and F1 maximization, highlighting consistent performance above 0.96 within the range of [0.55, 0.65].

C. Computational Performance and Failure Analysis

Detection exhibits $O(n)$ complexity averaging 127 ms (max 342 ms). Figure 3 shows the correlation ($R^2 = 0.386$, slope 0.245 ms/element) that confirms suitability in real-time. Memory averaged 12.3 MB (peak 23.7 MB).

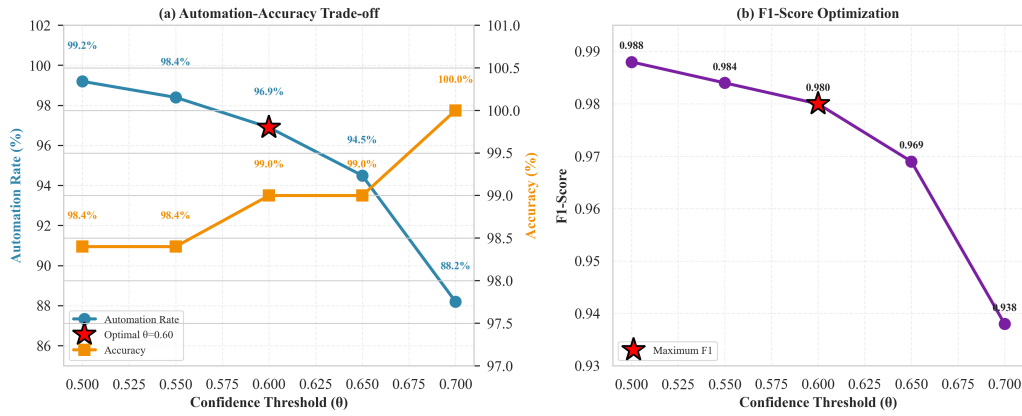


Fig. 2. Threshold optimization: (a) automation-accuracy trade-off; (b) F1-score stability analysis.

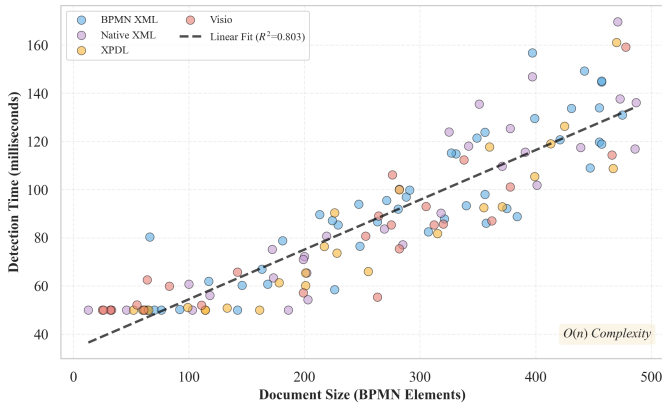


Fig. 3. Detection latency scaling with model complexity.

The single misclassification involved proprietary XML that mimicked the BPMN structure with vendor namespaces, yielding a marginal separation (0.68 versus 0.65) where $\Delta S = 0.03$ fell below $\epsilon = 0.15$ but exceeded $\theta = 0.6$. Distribution analysis reveals category separation: misclassified (0.665/0.030), ambiguous (0.385/0.025), correct (0.947/0.214), validating confidence-based quantification.

D. End-to-End Transformation and Comparative Evaluation

The pipeline evaluation measured integrated performance from detection through JSON standardization (Table V). The End-to-end transformation averaged 2.3 seconds with 98.7% semantic preservation, achieving 85% time reduction versus manual conversion. Format-specific precision ranged from 97.2% to 99.8%, with a unified JSON structure preserving the BPMN 2.0 taxonomy.

WSM achieves an optimal precision-latency balance with 99.0% detection at 127 ms versus AHP (inconsistent for $n > 4$ features), TOPSIS (82 ms overhead), and ELECTRE (non-compensatory). Table VI shows substantial advantages over the baseline. The suggested method achieves an impressive 99.0% accuracy coupled with 96.9% automation, yielding a 14.7% enhancement over existing benchmarks, while ensuring latency remains below 200 ms for real-time cross-platform transformation.

TABLE V
INTEGRATED TRANSFORMATION PIPELINE PERFORMANCE METRICS

Format	Detection	Parsing	E2E Time	Accuracy
BPMN XML	118 ms	1.7 s	1.8 s	99.8%
Native XML	134 ms	2.1 s	2.2 s	98.5%
XPDL	142 ms	2.3 s	2.4 s	98.1%
Visio	189 ms	2.9 s	3.1 s	97.2%
Average	127 ms	2.2 s	2.3 s	98.7%

TABLE VI
COMPARATIVE PERFORMANCE

Approach	Accuracy	Auto. Rate	Latency
Extension-only	84.3%	100%	23 ms
Namespace-only	87.4%	89.7%	67 ms
Manual selection	100%	0%	8.2 s
MCDA (proposed)	99.0%	96.9%	127 ms

E. Scalability and Batch Processing Considerations

The stateless detection algorithm enables parallel execution across document sets with linear memory scaling. The modular parser registry supports incremental format addition without algorithmic modifications, facilitating adaptation to emerging BPMN variants. Feature extraction operates on bounded memory windows with caching mechanisms that reduce redundant computation. The algorithm suits interactive workflows and automated pipelines, though further validation with larger datasets and corrupted inputs would strengthen deployment confidence.

V. CONCLUSION

This research addressed heterogeneous BPMN format identification through multi-criteria detection across file signatures, namespaces, structural patterns, and content heuristics. The Weighted Sum Model, optimized via grid search of 10,626 combinations with 5-fold cross-validation, resolves detection gaps using normalized scoring with adaptive thresholding ($\theta = 0.6$, $\epsilon = 0.15$).

Validation in 127 models demonstrated 99.0% accuracy with 96.9% automation, achieving a 2.3-second transformation with 98.7% semantic preservation through expert-validated mapping, representing 85% time reduction versus manual conversion. Compared to baselines (extension-only: 84.3%; namespace-only: 87.4%), the method achieves a 14.7% improvement while maintaining sub-200 ms latency. Ablation studies confirmed namespace superiority (11.8% decline) and dual-threshold effectiveness.

The framework is ready for pilot deployment in various areas, with a mechanism to effectively handle ambiguous inputs. Future prospects encompass adaptive weight adjustment, deep learning extraction, hierarchical categorization, behavioral verification, and streaming architecture. The modular architecture establishes an interoperable infrastructure that facilitates seamless BPMN transformation in various contexts while preserving semantic fidelity.

ACKNOWLEDGMENT

This work was funded by the Indonesian Endowment Fund for Education (LPDP) on behalf of the Indonesian Ministry of Higher Education, Science and Technology and managed under the EQUITY Program (Contract No 4299/B3/DT.03.08/2025 No 3029/PKS/ITS/2025); the Directorate General of Higher Education, Research and Technology of the Republic of Indonesia under Pendidikan Magister menuju Doktor untuk Sarjana Unggul (PMDSU) and Penelitian Sinergi Hilirisasi Inovasi Komersial; Institut Teknologi Sepuluh Nopember (ITS) under Riset Kolaborasi Indonesia, Penelitian Kemitraan A, and Penelitian Departemen; and Ethereum Foundation under Academic Grants Round 2025.

REFERENCES

- [1] T. Kräuter, A. Rutle, H. König, and Y. Lamo, "Formalization and analysis of bpmn using graph transformation systems," in *International Conference on Graph Transformation*. Springer, 2023, pp. 204–222.
- [2] K. R. Sungkono, R. Sarno, B. S. Onggo, and M. F. Haykal, "Enhancing model quality and scalability for mining business processes with invisible tasks in non-free choice," *Journal of King Saud University-Computer and Information Sciences*, vol. 35, no. 9, p. 101741, 2023.
- [3] M. Rosemann, J. v. Brocke, A. Van Looy, and F. Santoro, "Business process management in the age of ai—three essential drifts," *Information Systems and e-Business Management*, vol. 22, no. 3, pp. 415–429, 2024.
- [4] M. Abbasi, R. I. Nishat, C. Bond, J. B. Graham-Knight, P. Lasserre, Y. Lucet, and H. Najjaran, "A review of ai and machine learning contribution in business process management (process enhancement and process improvement approaches)," *Business Process Management Journal*, 2024.
- [5] B. Chandavarkar, S. Khatri *et al.*, "Heterogeneous data format integration and conversion (hdfic) using machine learning and ibm-dfdl for iot," *Evolving Systems*, vol. 15, no. 2, pp. 375–396, 2024.
- [6] J. Wang, Y. Feng, C. Shen, S. Rahman, and E. Kandogan, "Towards operationalizing heterogeneous data discovery," *arXiv preprint arXiv:2504.02059*, 2025.
- [7] D. N. Dolha and R. A. Buchmann, "Comparative analysis of natural language query responses on bpmn model serializations: Rdf graphs versus bpmn xml," in *International Conference on Informatics in Economy*. Springer, 2024, pp. 129–140.
- [8] K. R. Sungkono, R. Sarno, M. C. Salsabila, and C. P. Dewi, "A graph-based method for merging business process models by considering semantic similarity," *International Journal of Intelligent Engineering & Systems*, vol. 16, no. 2, 2023.
- [9] M. Rivadeh and S.-H. Mirian-Hosseiniabadi, "Formal translation of yawl workflow models to the alloy formal specifications: A testing application," *Software and Systems Modeling*, vol. 22, no. 3, pp. 941–968, 2023.
- [10] J. Park, M. Kim, and H. Lee, "Intelligent document format classification using machine learning for business process management," in *2024 International Conference on Advanced Information Networking and Applications*. Springer, 2024, pp. 245–256.
- [11] Q. Nivon and G. Salaün, "Automated generation of bpmn processes from textual requirements," in *International Conference on Service-Oriented Computing*. Springer, 2024, pp. 185–201.
- [12] T. Kampik, J. C. Nieves, and H. Lindgren, "Process model comparison and matching: a comprehensive survey," in *2023 IEEE International Conference on Process Mining*. IEEE, 2023, pp. 89–96.
- [13] G. A. Cutting and A.-F. Cutting-Decelle, "Intelligent document processing—methods and tools in the real world," *arXiv preprint arXiv:2112.14070*, 2021.
- [14] X. Zhang, B. Hu, S. Liu, Q. Sun, and L. Chen, "Attenflow: Context-aware architecture with consensus-based retrieval and graph attention for automated document processing," *Applied Sciences*, vol. 15, no. 13, p. 7517, 2025.
- [15] D. N. Dolha and R. A. Buchmann, "Generative ai for bpmn process analysis: experiments with multi-modal process representations," in *International Conference on Business Informatics Research*. Springer, 2024, pp. 19–35.
- [16] K. C. Febryanto, R. Sarno, A. F. Septiyanto, K. R. Sungkono, S. I. Sabilla *et al.*, "Leveraging graph-based for enhanced service identification and microservices partitioning in business process systems," in *2024 Beyond Technology Summit on Informatics International Conference (BTS-I2C)*. IEEE, 2024, pp. 119–124.
- [17] H. Zhang, X. Liu, Z. Wang, and J. Chen, "Automated transformation of workflow models: from yawl to bpmn," *Information Systems*, vol. 98, p. 101726, 2021.
- [18] M. N. Tiftik, T. G. Erdogan, and A. K. Tarhan, "A framework for multi-perspective process mining into a bpmn process model," *Mathematical Biosciences and Engineering*, vol. 19, no. 11, pp. 11 800–11 820, 2022.
- [19] Y. Peng, F. Bathelt, R. Gebler, R. Gött, A. Heidenreich, E. Henke, D. Kadioglu, S. Lorenz, A. Vengadeswaran, M. Sedlmayr *et al.*, "Use of metadata-driven approaches for data harmonization in the medical domain: scoping review," *JMIR Medical Informatics*, vol. 12, no. 1, p. e52967, 2024.
- [20] I. N. G. A. M. Wardhiana, M. Z. Ghuftron, S. M. Jati, A. F. Septiyanto, and R. Sarno, "Optimizing microservices architecture using multi-criteria decision making (mcdm): A case study of paypal rest api," in *2024 11th International Conference on Information Technology, Computer, and Electrical Engineering (ICITACEE)*. IEEE, 2024, pp. 308–313.
- [21] W. Wang, H. Hu, Z. Zhang, Z. Li, H. Shao, and D. Dahlmeier, "Document intelligence in the era of large language models: A survey," *arXiv preprint arXiv:2510.13366*, 2025.
- [22] L. Peña, D. Andrade, A. Delgado, and D. Calegari, "An approach for discovering inter-organizational collaborative business processes in bpmn 2.0," in *International Conference on Process Mining*. Springer, 2023, pp. 487–498.
- [23] G. Pal, K. Atkinson, and G. Li, "Managing heterogeneous data on a big data platform: A multi-criteria decision making model for data-intensive science," in *2020 IEEE International Conference on Big Data and Smart Computing (BigComp)*. IEEE, 2020, pp. 229–239.
- [24] A. S. Jadhav, "A novel weighted tpr-tnr measure to assess performance of the classifiers," *Expert systems with applications*, vol. 152, p. 113391, 2020.
- [25] M. Grandini, E. Bagli, and G. Visani, "Metrics for multi-class classification: an overview," *arXiv preprint arXiv:2008.05756*, 2020.
- [26] A. Rosa and A. Massaro, "Process mining organization (pmo) modeling and healthcare processes," *Knowledge*, vol. 3, no. 4, pp. 662–678, 2023.
- [27] G. Zhu, M. Ye, X. Yu, J. Liu, M. Wang, Z. Luo, H. Liang, and Y. Zhong, "Optimizing route planning via the weighted sum method and multi-criteria decision-making," *Mathematics*, vol. 13, no. 11, p. 1704, 2025.
- [28] S. Bandyopadhyay, "Comparison among multi-criteria decision analysis techniques: A novel method," *Progress in Artificial Intelligence*, vol. 10, no. 2, pp. 195–216, 2021.