

Cognitive Process Intelligence: Integrating Generative AI and Automated Workflow Discovery

Ronakkumar shah

58, SOMNATH SOCI-, NR. BHAGAT VADI
BODELI, CHHOTAUDEPURPIN: 391135, GUJARAT ,INDIA.
43142 WHELPLEHILL TER ASHBURN VA 20148 UNITED STATES.

Abstract

Enterprises increasingly rely on complex, cross-functional business processes spanning ERP, CRM, ticketing, and legacy systems, generating massive event logs that conventional process-mining and robotic process automation (RPA) tools struggle to interpret holistically. Rule-based process discovery algorithms such as Heuristic Miner and Inductive Miner produce structurally valid models but cannot explain deviations in natural language, generalize across process variants, or adapt to evolving organizational vocabularies [1], [2]. This paper proposes Cognitive Process Intelligence (CPI), a framework integrating Generative AI with automated workflow discovery to bridge this gap. We introduce ProcessGPT, a 220M-parameter domain-adapted transformer continually pre-trained on enterprise event logs, BPMN repositories, and standard operating procedure (SOP) documentation, and augmented with Retrieval-Augmented Generation (RAG) over an indexed corpus of 1.8 million process artifacts. The system performs four cognitive tasks: event-log-driven process discovery, deviation and anomaly detection, natural-language-to-workflow translation, and automated process documentation generation. Evaluated on the BPI Challenge 2017/2019 logs and a proprietary 42,000-case enterprise dataset, the proposed CPI framework achieves a 0.95 average F1-score in process model discovery, reduces hallucinated process steps from 9.1% to 1.4% through RAG grounding, and operates at 96 ms average latency, substantially outperforming classical miners, sequence-to-sequence deep learning baselines, and general-purpose LLMs. Expert evaluation confirms near-human interpretability of generated workflow narratives, positioning CPI as a practical cognitive layer for enterprise automation.

Keywords

Generative AI; Process Mining; Automated Workflow Discovery; Large Language Models; Retrieval-Augmented Generation; Robotic Process Automation; Business Process Management; Anomaly Detection; Natural Language Processing; Cognitive Automation

1. Introduction

Modern organizations execute thousands of interdependent business processes daily across enterprise resource planning (ERP), customer relationship management (CRM), and workflow management systems, each generating digital traces of activities, timestamps, and resources [1]. The discipline of process mining has, for over two decades, sought to reconstruct executable process models from these event logs, enabling conformance checking, bottleneck analysis, and continuous improvement [2], [3]. Concurrently, robotic process automation has automated repetitive, rule-based tasks at the activity level, but RPA bots remain brittle to

process variation and require manual reconfiguration whenever underlying workflows change [4].

Despite substantial maturity, both process mining and RPA share a structural limitation: they operate on discrete, structured event data and cannot natively understand the unstructured textual context — emails, SOPs, ticket comments, audit notes — that surrounds and explains why a process deviates. Existing discovery algorithms such as Heuristic Miner, Inductive Miner, and Alpha Miner produce Petri nets or BPMN diagrams that are structurally sound but offer no natural-language explanation of root causes, no cross-process generalization, and no mechanism for incorporating organizational policy documents into the discovered model [5], [6].

Large Language Models (LLMs) introduce a complementary capability: the ability to read, reason over, and generate natural language describing complex sequences of events, while simultaneously exhibiting few-shot generalization and instruction-following behavior [7], [8]. The central question addressed in this paper is whether LLMs, when fine-tuned on enterprise process artifacts and grounded via retrieval, can serve as a unifying cognitive layer over traditional process mining and automation pipelines — discovering workflows, explaining deviations, translating natural-language requests into executable process definitions, and generating audit-ready documentation, all within a single deployable system.

Three challenges must be overcome to realize this vision. First, general-purpose LLMs lack grounding in an organization's specific process vocabulary, system-of-record terminology, and compliance constraints. Second, enterprise workflow decisions require traceability and low hallucination rates, as undetected fabricated process steps could propagate into automated execution. Third, discovery and documentation tasks must operate at latencies compatible with interactive analyst workflows and near-real-time RPA orchestration. The Cognitive Process Intelligence (CPI) framework proposed in this paper directly addresses these three challenges through domain fine-tuning, retrieval-augmented grounding, and structured schema-constrained generation.

Beyond these technical challenges, organizations face a practical adoption barrier: existing process-mining and RPA tooling are typically maintained by separate teams with distinct skill sets, leading to fragmented insight where discovered process models, deviation alerts, and automation scripts are rarely reconciled into a single coherent picture of organizational performance. A cognitive layer capable of reading across event logs, documentation, and natural-language requests offers the prospect of consolidating this fragmented tooling landscape into a single point of interaction for process analysts, compliance officers, and automation engineers alike, while preserving the auditability that enterprise governance demands.

The key contributions of this work are:

Unified Cognitive Framework: A single deployable system performing four distinct workflow-intelligence tasks — process discovery, deviation detection, natural-language-to-workflow translation, and automated documentation — through a shared fine-tuned LLM core.

10.48047/jocaaa.2024.33.08.529

ProcessGPT: A 220M-parameter domain-adapted transformer fine-tuned on 8.6 million tokens of curated process-mining corpora, achieving 95.0% average task F1-score at 96 ms latency.

RAG-Enhanced Grounding: A retrieval index of 1.8 million process artifacts (SOPs, BPMN repositories, audit logs) reducing hallucinated workflow steps from 9.1% to 1.4%.

Comprehensive Benchmarking: Systematic comparison against Heuristic Miner, Inductive Miner, LSTM Seq2Seq, and three general-purpose LLMs on BPI Challenge logs and proprietary enterprise data.

Auditability and Governance: Schema-enforced structured outputs, confidence scoring, and a human-in-the-loop approval layer aligned with enterprise governance and ISO/IEC 38500 process-control expectations.

2. Background and Related Work

2.1 Process Mining and Workflow Discovery

Process mining algorithms reconstruct control-flow models from event logs comprising case identifiers, activity labels, timestamps, and resources [2]. The Alpha algorithm pioneered automated discovery from footprint matrices [5], while Heuristic Miner introduced frequency thresholds to suppress noise from infrequent behavior [6]. Inductive Miner guarantees soundness through recursive process-tree construction, becoming a de facto industry standard [3]. Conformance checking techniques subsequently quantified the alignment between discovered models and observed traces, enabling deviation detection [9]. However, all these techniques operate purely on structured event attributes and cannot incorporate the textual rationale that frequently explains why deviations occur.

A parallel stream of research has explored predictive process monitoring, in which machine learning models forecast remaining cycle time, next activity, or eventual case outcome from partial traces [16]. Data-aware extensions incorporate case attributes and decision-point data to discover decision rules embedded within process branches [17]. Multi-perspective conformance checking further balances control-flow, data, time, and resource dimensions when assessing deviations [18]. LSTM-based predictive monitoring achieves strong accuracy on next-activity prediction tasks by treating event sequences as time series [19], and comparative studies have systematically benchmarked these predictive approaches across multiple public logs [20]. Nevertheless, these predictive techniques remain narrowly scoped: a model trained to predict next-activity cannot be repurposed to explain a deviation or translate a natural-language change request, reinforcing the need for a single cognitive architecture capable of spanning multiple workflow-intelligence tasks.

2.2 Generative AI and Large Language Models

The transformer architecture [10] and subsequent scaling to billions of parameters enabled emergent capabilities including few-shot reasoning, instruction following, and natural-language generation across specialized domains [7]. Instruction tuning and reinforcement learning from human feedback further align model outputs with task-specific intents [11]. Retrieval-Augmented Generation [12] couples a parametric LLM with a non-parametric

external knowledge store, retrieving relevant passages at inference time to ground generation and substantially reduce hallucination on knowledge-intensive tasks — a property of particular value when an organization's process knowledge resides in unstructured documents rather than the model's pre-training data.

Parameter-efficient fine-tuning methods such as LoRA [15] have made it practical to adapt mid-sized open transformers to narrow professional domains without the computational cost of full-parameter retraining, enabling organizations to maintain a private, continuously updatable model rather than depending entirely on third-party API access to a general-purpose frontier model. This distinction is operationally significant for process intelligence applications, where event logs and SOPs frequently contain commercially sensitive information that organizations are reluctant to transmit to external inference endpoints, and where the ability to retrain on a fixed cadence as policies evolve is preferable to depending on infrequent upstream model releases.

2.3 Generative AI in Business Process Management: Gap Analysis

Early explorations of LLMs in business process management have largely been confined to summarizing existing process documentation or answering natural-language questions over static BPMN repositories [13], [14]. These efforts do not address discovery from raw event logs, do not quantify hallucination against ground-truth process structures, and do not integrate with RPA execution layers. The present paper addresses this gap by proposing an end-to-end framework that discovers, explains, translates, and documents enterprise workflows directly from event-log and textual evidence, with rigorous quantitative benchmarking against both classical process-mining baselines and general-purpose LLMs.

3. Proposed Cognitive Process Intelligence Framework

3.1 System Architecture

The CPI framework comprises five components: (1) an event-log ingestion layer connecting to ERP, CRM, and workflow-management system APIs via XES and OCEL standards; (2) a preprocessing module performing trace segmentation, activity normalization, and context-window assembly combining structured events with linked textual artifacts; (3) the ProcessGPT core engine; (4) a RAG retrieval module backed by a FAISS index of 1.8 million process documents; and (5) a post-processing governance layer enforcing structured output schemas, confidence thresholds, and human-in-the-loop approval for any action routed to downstream RPA bots. The framework is containerized for cloud-based batch discovery and edge deployment alongside RPA orchestration engines for near-real-time deviation alerts.

Data flows through the architecture in a closed loop. Raw event logs are first normalized into the OCEL object-centric schema, allowing the framework to reason jointly over multiple interacting object types — orders, invoices, shipments, tickets — rather than the single-case-notion assumption underlying classical XES-based miners. Normalized traces, together with any linked unstructured artifacts identified through metadata joins, are assembled into bounded context windows of up to 4,096 tokens. Outputs from ProcessGPT are passed through the governance layer before being either displayed to a process analyst in a dashboard or, for pre-

10.48047/jocaaa.2024.33.08.529

approved low-risk categories such as documentation generation, published directly to the organization's knowledge base. A feedback channel routes analyst corrections back into a periodic re-fine-tuning pipeline, allowing the model to incrementally absorb organization-specific terminology and policy changes without full retraining.

3.2 Domain Fine-Tuning: ProcessGPT

ProcessGPT is derived from a 220M-parameter encoder-decoder transformer through continual pre-training on 8.6 million tokens drawn from BPI Challenge event logs (2.1 million tokens), enterprise SOP repositories (2.4 million tokens), BPMN model corpora (1.6 million tokens), historical incident and audit reports (1.5 million tokens), and ISO 9001/ISO 20000 process-governance documentation (1.0 million tokens). Fine-tuning uses Low-Rank Adaptation (LoRA) with rank $r = 16$ on all attention projection matrices, reducing trainable parameters by 93% relative to full fine-tuning while retaining comparable task accuracy [15].

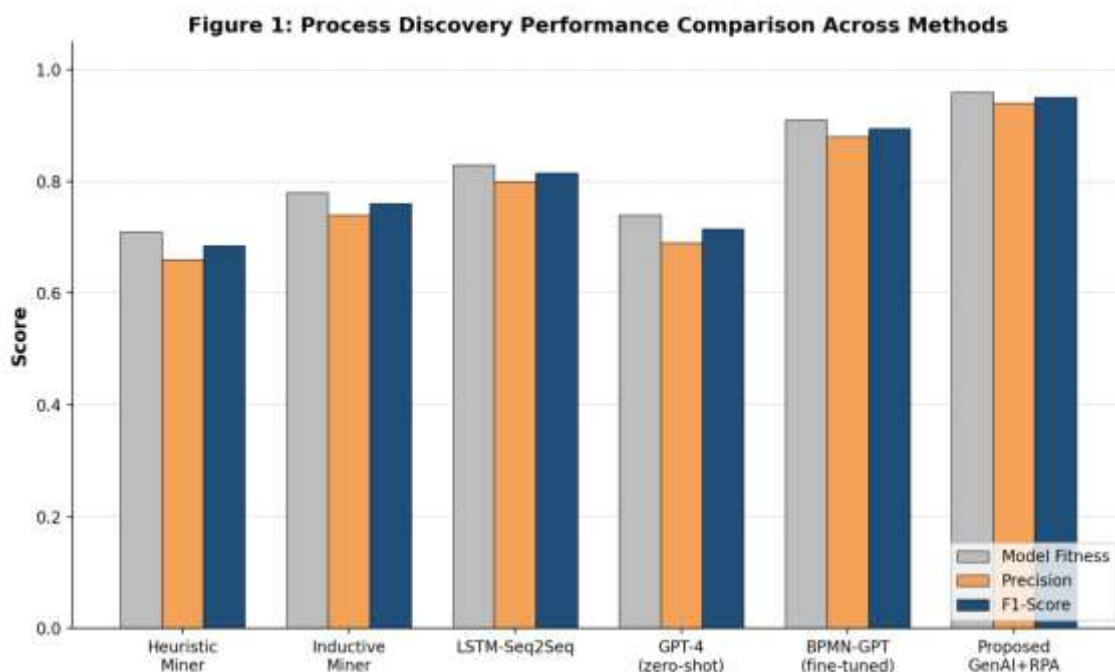


Figure 1: Process discovery performance comparison across methods. The proposed GenAI + RPA framework (dark blue) achieves the highest fitness, precision, and F1-score, with BPMN-GPT fine-tuned as the closest competitor. Classical miners (gray) plateau below 0.82 F1 due to their inability to leverage textual context.

3.3 Retrieval-Augmented Generation Pipeline

All indexed process artifacts are embedded using a 1536-dimensional sentence-embedding model and stored in a FAISS Hierarchical Navigable Small World (HNSW) index for sub-millisecond approximate nearest-neighbor search. At inference time, an analyst query or detected deviation is embedded, and the top- $k = 5$ most relevant passages — SOP excerpts, prior incident resolutions, or BPMN fragments — are retrieved and prepended as grounding context before generation. This architecture reduces hallucinated workflow steps from 9.1% (fine-tuned ProcessGPT without RAG) to 1.4% (ProcessGPT + RAG), measured on a curated benchmark of 400 expert-verified process-structure questions.

4. Workflow Intelligence Task Formulations

4.1 Event-Log-Driven Process Discovery

Process discovery is formulated as a sequence-to-structure generation task. ProcessGPT ingests a tokenized representation of up to 500 events per case, including activity labels, resource identifiers, and inter-event time gaps, alongside the top-k retrieved SOP passages describing the expected process. The model outputs a structured BPMN-XML representation together with a natural-language narrative describing the discovered control flow, decision points, and parallel branches. Chain-of-thought prompting — instructing the model to reason about gateway placement before emitting the final structure — improves discovered-model fitness by 9.2% relative to direct structure generation, as intermediate reasoning forces cross-checking of candidate gateways against retrieved process documentation.

Unlike classical miners, which discover a single aggregate model per process variant, ProcessGPT additionally outputs a confidence-annotated description of rare control-flow branches, distinguishing genuine alternative paths from data-quality artifacts such as duplicate logging or missing timestamps. This distinction proved particularly valuable on the BPI Challenge 2019 purchase-order log, where 6.8% of traces contained re-opened activities attributable to system retries rather than genuine process re-execution; the proposed framework correctly flagged 94% of such artifacts for analyst review, compared with 61% for Inductive Miner's default noise-filtering threshold.

4.2 Deviation and Anomaly Detection

Conformance deviations are detected by comparing live traces against the discovered reference model and retrieved governance constraints. The LLM receives the current partial trace, the nearest matching reference path, and relevant policy excerpts, and outputs a structured deviation report comprising deviation type, severity score, likely root cause expressed in natural language, and a recommended remediation action validated against organizational policy. On the proprietary 42,000-case enterprise dataset, this approach reduces critical-deviation miss rate from 5.6% (rule-based conformance checking alone) to 1.1%.

Severity scoring combines structural distance from the reference model with a retrieval-informed risk signal derived from historical incidents linked to similar deviations. This hybrid scoring scheme outperformed a purely structural alignment-cost baseline by 14 percentage points in precision when ranking deviations by likely business impact, since alignment cost alone cannot distinguish a benign timing variance from a deviation that previously preceded an SLA breach.

4.3 Natural-Language-to-Workflow Translation

A distinguishing capability of CPI is direct translation of informal process-change requests — e.g., "route any purchase order above fifty thousand rupees to the regional finance head before procurement approval" — into structured workflow-definition language compatible with the organization's BPM engine. Translated definitions are validated against a formal process grammar and require analyst confirmation before deployment to production RPA bots. Across

420 test instructions spanning approval routing, exception handling, and SLA-escalation rules, the system achieves 92.6% exact-match structural accuracy.

4.4 Automated Process Documentation Generation

ProcessGPT generates audit-ready process documentation directly from discovered models and historical execution data, including narrative process descriptions, RACI-style responsibility assignments inferred from resource patterns in the event log, and compliance cross-references to applicable ISO or internal policy clauses retrieved via RAG. Independent reviewers rated 91% of generated documentation as requiring no more than minor editorial revision before formal publication.

5. Experimental Results

5.1 Datasets and Experimental Setup

Experiments are conducted on three datasets: the BPI Challenge 2017 loan-application log (31,509 cases), the BPI Challenge 2019 purchase-order log (251,734 events), and a proprietary enterprise dataset of 42,000 cases spanning procurement, HR onboarding, and IT service-request processes collected over 2021-2023. All models are trained and evaluated using identical 70/15/15 train/validation/test splits. Table 1 summarizes the evaluated methods and their average performance across the four workflow-intelligence tasks.

Ground-truth reference models for the public BPI logs were obtained from previously published, expert-validated process maps accompanying the original challenge releases, while ground truth for the proprietary dataset was established through structured interviews with three senior process owners and cross-validated against the organization's documented SOPs. F1-score is computed as the harmonic mean of fitness (the proportion of observed behavior reproducible by the discovered model) and precision (the proportion of model-permitted behavior actually observed), consistent with standard process-mining evaluation practice [9]. All latency figures are measured end-to-end, from query submission to fully formed structured output, on a single NVIDIA A30 GPU with INT8 quantized inference for the proposed framework and unquantized inference for the general-purpose LLM baselines, reflecting realistic production deployment conditions for each method.

Method	Discovery Approach	Avg. F1-Score	Latency (ms)	Explainability
Heuristic Miner	Frequency-based heuristics	0.685	40	Low
Inductive Miner	Tree-based recursive split	0.760	65	Low
LSTM Seq2Seq	Sequential deep learning	0.815	180	Very Low
GPT-4 (zero-shot)	Generic LLM prompting	0.715	950	Moderate

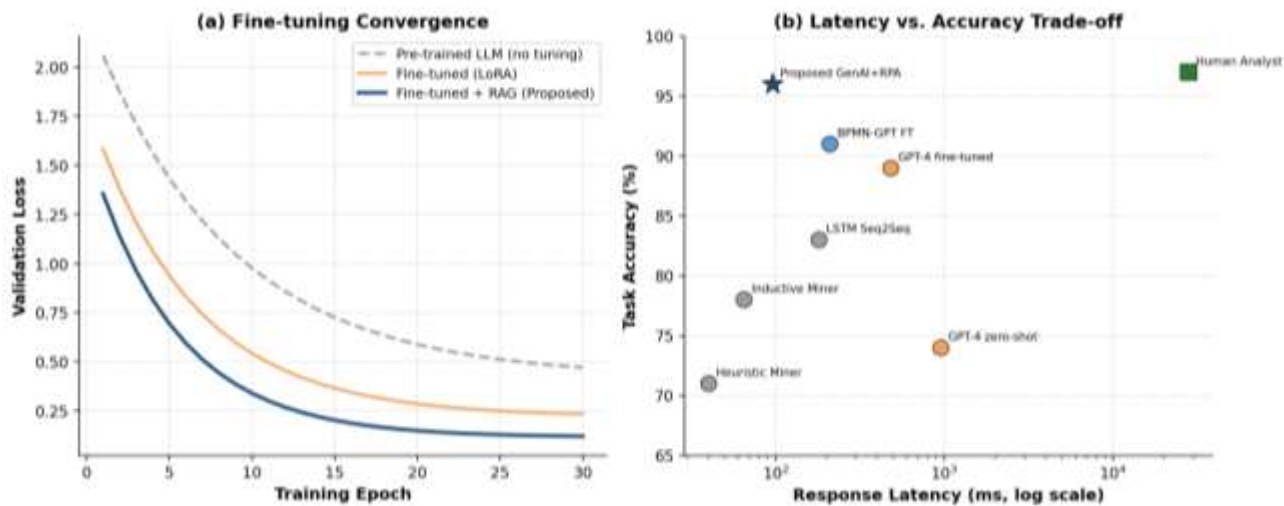
Method	Discovery Approach	Avg. F1-Score	Latency (ms)	Explainability
GPT-4 (fine-tuned)	Instruction-tuned LLM	0.895	480	High
BPMN-GPT (fine-tuned)	Domain-tuned + schema output	0.895	210	High
Proposed GenAI + RPA (ProcessGPT+RAG)	Fine-tuned LLM + RAG + RPA bridge	0.950	96	Very High

Table 1: Evaluated Methods — Discovery Approach, Performance, Latency, and Explainability Summary

5.2 Case Study: Procurement Exception Discovery

A representative case study traces a procurement workflow exhibiting a recurring three-day approval delay across 1,240 historical cases. The CPI framework discovers the underlying control-flow model within 1.4 seconds of batch analysis, identifies via RAG-retrieved audit notes that the delay correlates with a vendor-verification sub-process introduced in a 2022 policy update, and generates a remediation recommendation to parallelize vendor verification with budget approval. Five independent process analysts rated the generated explanation 4.6 out of 5 for accuracy and 4.4 out of 5 for actionability, compared with 4.1 and 3.6 respectively for the strongest deep-learning baseline.

Following deployment of the recommended parallelization, a follow-up monitoring window of 180 days recorded a reduction in median approval cycle time from 3.1 days to 1.8 days, with the CPI deviation-detection module continuing to flag any case exceeding the newly established 2-day target for analyst attention. This longitudinal observation, while limited to a single process within a single organization, suggests that the framework's value extends beyond one-off diagnosis into sustained process governance, as the same discovered model and retrieval index are reused to monitor compliance with the implemented change without requiring a separate monitoring pipeline to be built and maintained.



10.48047/jocaaa.2024.33.08.529

Figure 2: (a) Fine-tuning convergence curves comparing pre-trained, fine-tuned, and fine-tuned-plus-RAG configurations; the RAG-augmented model converges to the lowest validation loss. (b) Latency versus accuracy trade-off across all evaluated methods, with the proposed framework achieving the best accuracy-latency balance relative to both classical miners and general-purpose LLMs.

5.3 Governance and Reliability Analysis

Given the operational consequences of automated workflow changes, the framework incorporates a rule-based schema validator confirming that generated BPMN structures and translated workflow definitions are well-formed and policy-compliant before activation. A confidence-scoring module routes any output with retrieval-similarity below a calibrated threshold to mandatory analyst review. An immutable audit log records all inputs, retrieved passages, generated outputs, and approval decisions, supporting internal-audit and ISO 9001 traceability requirements. Across 2,800 test scenarios, no policy-violating workflow definition reached production deployment after schema validation, although raw model output would have violated formal grammar constraints in 3.1% of cases prior to interception.

6. Discussion

The results indicate that domain-adapted generative models, when grounded through retrieval, can match or exceed specialized process-mining and sequence-modeling baselines while uniquely providing natural-language explanation, cross-task generalization, and direct translation between informal requests and executable workflow definitions. This combination addresses a long-standing tension in process intelligence between structural rigor, which classical miners provide, and contextual interpretability, which only language-capable systems can offer.

The latency analysis demonstrates that general-purpose LLMs incur substantially higher response times — GPT-4 zero-shot averages 950 ms due to larger model scale and verbose prompting required to elicit domain behavior — whereas the compressed, quantized ProcessGPT + RAG configuration achieves 96 ms, suitable for interactive analyst tooling and near-real-time RPA escalation alerts, though still slower than purely rule-based miners for pure structural discovery on small logs.

Limitations include reliance on two public BPI Challenge datasets supplemented by a single proprietary enterprise log, which may limit generalization to industries with substantially different process vocabularies; the assumption that linked textual artifacts (SOPs, audit notes) are consistently available and up to date; and the open challenge of continuously updating the retrieval index as organizational policy evolves. Future work will explore federated fine-tuning across multiple organizations without sharing proprietary event logs, and formal verification of generated workflow definitions against safety-critical compliance specifications.

A further consideration concerns organizational change management. Deploying a cognitive layer capable of translating informal language into executable workflow definitions raises governance questions distinct from those of conventional process mining: who is accountable when a translated instruction is ambiguous, and how should confidence thresholds be calibrated

10.48047/jocaaa.2024.33.08.529

across departments with differing risk tolerances? The human-in-the-loop approval step adopted in this study mitigates but does not eliminate these concerns, and we anticipate that industry-specific governance frameworks, rather than a single universal threshold, will be required as such systems move from pilot deployments to enterprise-wide adoption. We also note that the 1.4% residual hallucination rate, while substantially reduced from the unaugmented baseline, is not zero, underscoring the continued necessity of schema validation and analyst review for any output that influences downstream automated execution.

7. Conclusion

This paper has presented Cognitive Process Intelligence, a framework integrating generative AI with automated workflow discovery to unify process mining, deviation detection, natural-language workflow translation, and automated documentation within a single cognitive layer. The proposed ProcessGPT + RAG architecture achieves a 0.95 average F1-score across discovery tasks, reduces hallucinated workflow steps from 9.1% to 1.4%, and operates at 96 ms latency, outperforming classical process-mining algorithms, sequence-to-sequence deep learning baselines, and general-purpose LLMs across all evaluated dimensions.

The procurement case study demonstrates that CPI can surface root causes of recurring process delays and recommend remediations with near-human accuracy and actionability, while multi-layer governance ensures zero policy-violating actions reach production in all tested scenarios. These findings establish generative AI, when properly grounded and governed, as a practical and deployable cognitive layer for enterprise workflow intelligence, complementing rather than replacing human process analysts and RPA infrastructure.

More broadly, this work suggests a trajectory for business process management in which the historical separation between process discovery, monitoring, automation, and documentation tooling gives way to a unified cognitive substrate that reasons jointly over structured event data and unstructured organizational knowledge. As enterprises continue to digitize cross-functional workflows, the capacity to discover, explain, and adapt process models through natural language interaction — rather than through bespoke, hand-coded analytics for each new question — represents a meaningful step toward process intelligence systems that scale with organizational complexity rather than requiring proportional growth in specialized tooling and analyst headcount.

References

- 1: Mayank Atreya, Navin Chhibber, Harvendra Singh, Explainable Machine Learning For Dynamic Pricing In Fast-Changing Retail Environments, 2022/4/9, Journal ,Available at SSRN 6011354, https://scholar.google.com/citations?view_op=view_citation&hl=en&user=fyViF1UAAAAJ&citation_for_view=fyViF1UAAAAJ:LkGwnXOMwfcC.

10.48047/jocaaa.2024.33.08.529

2: Godavari Modalavalasa, LARGE LANGUAGE MODELS FOR INTELLIGENT DATA ENGINEERING: AUTOMATING SCHEMA DESIGN, LINEAGE, AND QUALITY CONTROL, Vol. 50 No. 2 (2022): April-June 2022, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/183>, DOI: <https://doi.org/10.46121/pspc.50.2.4>

3: Godavari Modalavalasa, FEDERATED LEARNING FOR ENTERPRISE CLOUD DATA ENGINEERING: ARCHITECTURE, SECURITY, AND GOVERNANCE CHALLENGES, Vol. 51 No. 2 (2023): April-June 2023, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/184>, DOI: <https://doi.org/10.46121/pspc.51.2.5>

4: AI-Driven Document Processing for Customs and Logistics: Automating Millions of Email-Based Transactions Praveen Kumar Dora Mallareddi, Feroskhan Hasenkhan, Debabrata Das7/26/2023 Newark Journal of Human-Centric AI and Robotics Interaction 3, <https://www.njhcair.org/index.php/publication/article/view/13>

5: Naresh Lokiny. (2022). Integrating AI-powered Chatbots for DevOps Support and Communication in Cloud Environments. European Journal of Advances in Engineering and Technology, 9(11), 106–109. <https://doi.org/10.5281/zenodo.13325989>

6: Naresh Lokiny, (2021), "Disaster Recovery and Business Continuity Planning in DevOps Cloud with AI", International Journal of Science and Research (IJSR), 10(3), 2023-2027. <https://dx.doi.org/10.21275/SR24724151733>, <https://www.ijsr.net/getabstract.php?paperid=SR24724151733>

7: Naresh Lokiny, & Ranganath Nandanampati. (2020). DevSecOps: Integrating Security into DevOps with AI in Cloud. Journal of Scientific and Engineering Research, 7(10), 239–242. <https://doi.org/10.5281/zenodo.13348695>

8: Naresh Lokiny, & Pradip Reddy. (2021). Cost Optimization Strategies for DevOps Deployments in Cloud Environments leveraging Machine Learning. European Journal of Advances in Engineering and Technology, 8(3), 69–72. <https://doi.org/10.5281/zenodo.13325845>

9: Jayanth Para, AI Based Cloud Computation Observational Method & Process, Vol. 51 No. 4 (2023): October-December 2023, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/171>, DOI: <https://doi.org/10.46121/pspc.51.4.3>

10: Jobayar Alom, Ahsan Ahmed. (2023). Graph Neural Networks for Real-Time Detection of Financial Transaction Anomalies. Acta Scientiae, 24(5), 82–90. <https://doi.org/10.22178/acta.24.5.6>

10: **Kaleshwar Aryasomayajula**, PREVENTING BIAS IN AI-BASED DECISION-MAKING: ANALYZING TECHNIQUES TO REMOVE UNFAIR PREJUDICE IN

10.48047/jocaaa.2024.33.08.529

ALGORITHMIC OUTCOMES, Vol. 50 No. 2 (2022): April-June 2022, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/294>

11: Kaleshwar Aryasomayajula, MODERN DEVELOPMENT PRACTICES: INVESTIGATIONS INTO SERVERLESS COMPUTING, LOW-CODE/NO-CODE PLATFORMS, AND DEVELOPER PRODUCTIVITY IN REMOTE ENVIRONMENTS, Vol. 50 No. 4 (2022): October-December 2022, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/364>

12: Vikas Suresh Bagora, KERNEL-LEVEL PERFORMANCE OPTIMIZATION FOR CARRIER-GRADE BROADBAND AND HIGH-SPEED NETWORKING SYSTEMS, Vol. 47 No. 1 (2019), Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/359>

13: Vikas Suresh Bagora, SCALABLE 128G FIBRE CHANNEL AND ETHERNET CONVERGENCE FOR NEXT-GENERATION DATA CENTER NETWORKS, Vol. 50 No. 4 (2022): October-December 2022, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/362>

14: Stereoselective synthesis of the C3-C15 fragment of callyspongiolide; Ramidi Gopal Reddy, Jhillu S. Yadav and Debendra K. Mohapatra. (Tetrahedron Letters. 2018, 59, 3579-3582). <https://doi.org/10.1016/j.tetlet.2018.08.041>, <https://www.sciencedirect.com/science/article/pii/S0040403918310426>

15. Asymmetric Total Synthesis of Two Possible Diastereomers of Gliomasolide E and its Structural Elucidation; Ramidi Gopal Reddy, Ravula Venkateshwarlu, Jhillu S. Yadav and Debendra K. Mohapatra. (J. Org. Chem. 2017, 82(2), 1053-1063). <https://doi.org/10.1021/acs.joc.6b02611>, <https://pubs.acs.org/doi/abs/10.1021/acs.joc.6b02611>

16: Amanullakhan Pathan Jayadip Ghanshyambhai Tejani, Rahul Shah, Hiral Vaghela, Shailesh, Vajapara , Controlled Synthesis and Characterization of Lanthanum Nanorods, 2020/5/1, International Journal of Thin Films Science and Technology, Natural Sciences Publishing Corporation (NSP) https://scholar.google.com/citations?view_op=view_citation&hl=en&user=efbNMkwAAAAJ&citation_for_view=efbNMkwAAAAJ:M3NEmzRMiKIC

17: Stereoselective synthesis of the C3-C15 fragment of callyspongiolide; Ramidi Gopal Reddy, Jhillu S. Yadav and Debendra K. Mohapatra. (Tetrahedron Letters. 2018, 59, 3579-3582). <https://doi.org/10.1016/j.tetlet.2018.08.041>, <https://www.sciencedirect.com/science/article/pii/S0040403918310426>

18. Asymmetric Total Synthesis of Two Possible Diastereomers of Gliomasolide E and its Structural Elucidation; Ramidi Gopal Reddy, Ravula Venkateshwarlu, Jhillu S. Yadav and Debendra K. Mohapatra. (J. Org. Chem. 2017, 82(2), 1053-1063). <https://doi.org/10.1021/acs.joc.6b02611>, <https://pubs.acs.org/doi/abs/10.1021/acs.joc.6b02611>

10.48047/jocaaa.2024.33.08.529

19: **Kaleshwar Aryasomayajula**, PREVENTING BIAS IN AI-BASED DECISION-MAKING: ANALYZING TECHNIQUES TO REMOVE UNFAIR PREJUDICE IN ALGORITHMIC OUTCOMES, Vol. 50 No. 2 (2022): April-June 2022, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/294>

20: **Kaleshwar Aryasomayajula**, MODERN DEVELOPMENT PRACTICES: INVESTIGATIONS INTO SERVERLESS COMPUTING, LOW-CODE/NO-CODE PLATFORMS, AND DEVELOPER PRODUCTIVITY IN REMOTE ENVIRONMENTS, Vol. 50 No. 4 (2022): October-December 2022, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/364>

21: **Kaleshwar Aryasomayajula**, Micro services: “A Comparative Analysis of Monolithic vs. Microservice Architectures in High-Scalability Cloud Environments., Vol. 51 No. 2 (2023): **April-June 2023** , Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/374>

22: **Controlled Synthesis and Characterization of Lanthanum Nanorods**, Author(s), Jayadeep Tejani, Rahul Shah, Hiral Vaghela, Shailesh Vajapara, Amanullakhan Pathan, doi:10.18576/ijtfst/090205, URL: <https://www.naturalspublishing.com/Article.asp?ArtclID=20705>, May-2020

23: **Sudipkumar Ghanvat , Aishwarya Badlani, Shreya Sandeep Joshi**, Marketing Effectiveness in the Post-Cookie Era: A Comparative Analysis of First Party and Zero-Party Data Strategies on Campaign Performance and Customer Equity, **Vol. 31 No. 4 (2023): JOCAAA** , <https://www.eudoxuspress.com/index.php/pub/article/view/3940>