



Hybrid LLM and Knowledge Graph Framework for Explainable Healthcare Claims Adjudication

Selvakumar Kalyanasundaram
Independent Researcher, Texas, USA.

Received On: 23/06/2026

Revised On: 25/07/2026

Accepted On: 03/08/2026

Published On: 10/08/2026

Abstract - Healthcare claims adjudication remains one of the most complex, high-stakes, and error-prone processes in modern health insurance operations. Manual review is costly, inconsistent, and inherently opaque. Automated rule-based engines improve throughput but lack adaptability and generate minimal audit-ready reasoning. This paper presents KLARA (Knowledge-augmented LLM Adjudication and Reasoning Architecture), a novel hybrid framework that tightly integrates large language models (LLMs) with a medical knowledge graph (MedKG) to adjudicate insurance claims with high accuracy, clinical fidelity, and full explainability. KLARA encodes clinical coding standards (ICD-11, CPT-4, HCPCS Level II), payer-specific benefit policies, and clinical evidence pathways as a richly interconnected graph. At inference time, a retrieval-augmented generation (RAG) pipeline grounds the LLM's adjudication decisions in verified, authoritative graph nodes, preventing hallucination while preserving the model's natural-language reasoning capacity. A chain-of-thought (CoT) prompting strategy elicits structured, step-by-step justification for every adjudication decision. Extensive experiments on three benchmarked datasets CMS synthetic claims, a de-identified commercial payer corpus ($N = 124,000$), and a publicly released Medicare Part B sample demonstrate that KLARA achieves 94.7% overall adjudication accuracy, a 38% reduction in false denial rate relative to the best competing baseline, and an explainability fidelity score of 0.91 on the SHAP-aligned XAI metric. Clinician reviewers rated 89.4% of KLARA's natural-language justifications as "clinically sufficient." These results establish KLARA as a significant advance toward trustworthy, auditable AI in healthcare revenue cycle management.

Keywords - Claims Adjudication, Explainable Artificial Intelligence (XAI), and Healthcare Informatics, Knowledge Graphs, Large Language Models (Llms), Retrieval-Augmented Generation (RAG), and Revenue Cycle Management.

1. Introduction

Healthcare claims adjudication is the process by which a payer insurer, government program, or self-insured employer determines whether a submitted claim is reimbursable, the appropriate payment amount, and the clinical and contractual basis for that determination. The United States healthcare system alone processes over 4.5 billion insurance claims annually, representing more than \$4.3 trillion in submitted charges [1]. Despite the economic magnitude of this process, it remains heavily reliant on rigid deterministic rule engines that encode policy logic as static boolean trees and on human reviewers who impose subjective judgment under cognitive load.

The consequences of adjudication error are severe. Improper denial of medically necessary care creates access barriers and contributes to patient mortality in time-sensitive diagnoses such as oncological intervention and cardiac procedures [2]. Erroneous approval of fraudulent or unbundled claims drives an estimated \$60 billion in Medicare and Medicaid improper payments annually [3]. Critically, conventional systems provide minimal traceability: when a claim is denied, the rationale is typically a three-character remark code that offers the provider no actionable path to correction, fueling a costly appeals

ecosystem that consumes an estimated 34.2 billion administrative hours per year industry-wide [4].

Large language models (LLMs) present a compelling opportunity to transform this landscape. Models trained on vast clinical, regulatory, and administrative corpora demonstrate sophisticated language understanding, structured reasoning over complex multi-constraint problems, and the capacity to generate natural-language rationale for their decisions. However, LLMs applied naively to claims adjudication introduce unacceptable risks: hallucination of clinical codes or benefit rules that do not exist; inability to consistently enforce payer-specific contractual policy; and lack of deterministic auditability required by CMS, HIPAA, and state insurance regulatory frameworks [5], [6].

Knowledge graphs (KGs) offer a complementary capability profile. A well-constructed medical knowledge graph can encode authoritative clinical coding taxonomies, policy rule logic, evidence-based clinical pathways, and inter-entity relationships with formal semantics. KG traversal is deterministic, auditable, and immune to the hallucination failure modes of generative models. However, KGs lack the flexible language understanding and contextual inference required to interpret ambiguously documented clinical

scenarios, extract intent from free-text clinical notes, or generate human-readable explanations calibrated to the recipient [7].

This paper presents KLARA (Knowledge-augmented LLM Adjudication and Reasoning Architecture), a hybrid system that integrates LLM reasoning with knowledge graph grounding through a retrieval-augmented generation (RAG) pipeline. KLARA is designed around three core desiderata: (i) accuracy on par with best-in-class rule engines while handling clinical complexity that defeats rigid rules; (ii) full explainability through structured, citation-linked natural-language justifications; and (iii) compliance auditability meeting CMS and HIPAA requirements for algorithmic decision transparency.

The remainder of this paper is organized as follows. Section II surveys related work on AI in claims processing and XAI in healthcare. Section III details the KLARA architecture, including MedKG construction, RAG pipeline design, and CoT prompting strategy. Section IV describes experimental datasets, evaluation metrics, and baselines. Section V presents results and ablation studies. Section VI provides discussion including limitations and ethical considerations. Section VII concludes the paper.

2. Related Work

2.1. Automated Claims Processing and Rule Engines

Legacy claims adjudication platforms such as TriZetto Facets, Epic Tapestry, and Cognizant TriZetto rely on decision tables and Drools-based rule engines that encode payer benefit logic as structured IF-THEN predicates. Bhatt et al. [8] characterized the scalability limitations of such systems, noting that enterprise rule bases can contain over 250,000 individual rules, with cross-rule dependencies creating combinatorial complexity that resists systematic testing and maintenance. Johnson & Patel [9] demonstrated that rule engine error rates increase nonlinearly with policy complexity, particularly for claims involving comorbid conditions requiring multi-step medical necessity evaluation.

Machine learning approaches to claims automation emerged in the mid-2010s, primarily targeting fraud detection [10], [11] and denial prediction [12]. Gradient boosting models trained on structured claim features (procedure code, diagnosis, provider specialty, place of service) achieved AUCs of 0.82-0.88 on denial prediction but could not directly adjudicate claims or provide interpretable denial reasoning [13]. Neural approaches using CPT code embeddings showed improvement in capturing procedural similarity [14] but retained the same interpretability deficit.

2.2. Large Language Models in Clinical and Administrative Healthcare

The clinical NLP literature has evolved rapidly since the emergence of transformer-based models. BioBERT [15], ClinicalBERT [16], and GatorTron [17] demonstrated strong performance on entity recognition, relation extraction, and clinical question answering but were not applied to

administrative claims contexts. Med-PaLM 2 [18] and GPT-4 with clinical prompting [19] exhibited expert-level performance on clinical reasoning benchmarks including USMLE and MedQA, motivating application to complex insurance medical necessity review.

Studies directly applying LLMs to claims processing are nascent. Chen et al. [20] demonstrated that GPT-4 could correctly classify 78.3% of prior authorization requests when provided benefit criteria in context but noted significant variance attributable to hallucinated policy citations. Nakamura et al. [21] applied instruction-tuned LLaMA variants to claim denial letter generation and found high fluency but low factual grounding, with 23% of generated rationales containing at least one incorrect code reference. These findings motivate our KG-grounded approach.

2.3. Knowledge Graphs in Healthcare

Medical knowledge graphs have been extensively developed for clinical decision support. UMLS [22], SNOMED CT [23], and the National Drug File-Reference Terminology (NDF-RT) provide canonical entity representations. General-purpose biomedical KGs including PrimeKG [24] and DRKG [25] integrate multi-source evidence for drug-gene-disease relationships. The Open Biological and Biomedical Ontology (OBO) Foundry [26] provides formal ontological scaffolding. However, none of these resources specifically encodes payer-side adjudication policy, CPT/HCPCS billing relationships, or benefit plan structures, which are essential to claims processing.

KG-augmented reasoning has demonstrated strong results in clinical decision support [27], drug-drug interaction detection [28], and rare disease diagnosis [29]. Retrieval-augmented generation over KGs was formalized by Lewis et al. [30] and has been adapted for biomedical QA by Yasunaga et al. [31] (QA-GNN) and Luo et al. [32] (DRAGON). These works inform our RAG-CoT pipeline but do not address the administrative claims domain.

2.4. Explainability in Healthcare AI

Regulatory and ethical frameworks increasingly mandate explainability for AI systems in high-stakes domains [33]. In healthcare, the EU AI Act classifies medical AI as ‘high-risk,’ requiring human oversight and technical documentation of decision logic. In the United States, the CMS Interoperability and Patient Access final rule and ONC Health IT certification criteria impose transparency requirements on algorithmic payer decisions [34]. Post-hoc XAI methods including SHAP [35], LIME [36], and integrated gradients [37] have been applied to clinical predictive models but generate feature attributions that are difficult to communicate to clinical or administrative stakeholders. Our work instead pursues ante-hoc explainability through structured chain-of-thought reasoning grounded in authoritative KG evidence.

3. Methodology: The Clara Architecture

KLARA operates as a five-stage pipeline: (1) claim ingestion and normalization; (2) MedKG query and subgraph

extraction; (3) RAG-augmented context assembly; (4) LLM adjudication with CoT prompting; and (5) structured decision output with audit trail. Figure 1 provides a schematic overview.

3.1. MedKG: Medical Knowledge Graph Construction

The foundational component of KLARA is MedKG, a purpose-built knowledge graph integrating clinical, administrative, and policy knowledge sources. MedKG comprises 2.4 million nodes and 18.7 million edges across seven node types and fourteen edge relation types.

- **Node Taxonomy:** MedKG encodes the following entity classes: (a) Diagnosis entities from ICD-11 (80,230 billable codes with hierarchical parent-child relationships and inclusion/exclusion notes); (b) Procedure entities from CPT-4 (11,079 codes with RVU weights, bundling rules, and laterality modifiers); (c) Supply and device entities from HCPCS Level II (6,872 codes); (d) Drug entities from RxNorm (104,000 concept unique identifiers with ingredient, form, and strength attributes); (e) Benefit policy entities encoding payer-specific coverage determinations (local and national coverage determinations from CMS, 847 Medicare NCDs and 4,200+ LCDs, and 12 commercial payer policy libraries); (f) Clinical evidence entities linking procedures and diagnoses to USPSTF recommendations, NCCN guidelines, and AHA/ACC pathway nodes; (g) Provider entities encoding specialty, NPI taxonomy, and CMS certification status.
- **Edge Semantics:** Fourteen relation types are defined with formal OWL-compatible semantics, including: IS_INDICATED_FOR (diagnosis → procedure); REQUIRES_PRIOR_AUTH (procedure → payer-policy); BUNDLES_WITH (CPT → CPT, per CCI edits); EXCLUDES (ICD-11 exclusion notes); HAS_EVIDENCE_LEVEL (procedure → clinical-guideline); COVERED_UNDER (service → benefit-plan); and MEDICALLY_NECESSARY_IF (procedure → diagnostic-criterion). Edges carry provenance metadata including source document, effective date, and confidence weight derived from evidence grading.
- **Graph Maintenance:** MedKG is maintained through a continuous ingestion pipeline that monitors CMS MLN Matters, AMA CPT editorial updates, LCD/NCD revisions, and payer policy bulletins. Versioning is implemented at the node and edge level using named graph semantics in the underlying RDF triplestore (Apache Jena TDB2), enabling historical adjudication replay for audit purposes.

3.2. Claim Normalization and Entity Linking

Submitted claims arrive in ANSI X12 837P/837I format and are parsed into a structured claim object comprising: subscriber demographics, plan identifiers, rendering provider NPI, place of service, service date range, diagnosis codes (ICD-11, up to 12), procedure codes (CPT/HCPCS, up to 50

line items), revenue codes for institutional claims, and any attached clinical documentation in CCD-A or PDF format.

Entity linking maps submitted codes to MedKG node IDs. For ICD-11 and CPT codes, mapping is direct via identifier lookup. For free-text fields (clinical documentation, prior authorization narratives), a fine-tuned BioLinkBERT-large model [38] performs named entity recognition and UMLS-grounded entity disambiguation, resolving clinical mentions to MedKG concept nodes with confidence scores. Entities with link confidence below 0.72 are flagged for human review rather than passed to automated adjudication.

3.3. Subgraph Extraction and RAG Context Assembly

Given the linked claim entities, KLARA executes a multi-hop graph traversal to extract a claim-specific subgraph S_x . The traversal follows a priority-ordered exploration policy: (1) direct edges between submitted diagnosis and procedure nodes; (2) CCI bundling and unbundling rules applicable to the procedure set; (3) applicable benefit policy nodes for the identified plan-procedure pairs; (4) medical necessity criteria nodes for procedures flagged as requiring authorization; (5) evidence-graded guideline nodes supporting or contradicting the clinical scenario.

Subgraph S_x is serialized as a structured context document using a schema-directed linearization algorithm. Each node is rendered as a titled fact block with its key attributes; each edge is rendered as a declarative relationship sentence. The serialized subgraph context is bounded at 12,000 tokens to accommodate LLM context window constraints, with a priority-based truncation policy that preserves benefit policy and medical necessity nodes when truncation is necessary.

Vector-based retrieval supplements graph traversal for policy edge cases. A policy library of 18,400 claim scenario vignettes (curated from historical adjudications) is embedded using text-embedding-3-large and stored in a FAISS index. The top-k=5 semantically similar vignettes are retrieved and appended to the RAG context, providing few-shot analogical grounding for novel claim configurations not fully covered by the graph.

3.4. LLM Adjudication with Chain-of-Thought Prompting

KLARA employs GPT-4o as the primary reasoning model, accessed via the Azure OpenAI Service within a HIPAA Business Associate Agreement. The prompt architecture consists of four components:

- **System Prompt:** Establishes the model's role as a certified claims adjudicator, specifies the decision taxonomy (APPROVE / DENY / PEND-FOR-INFO / PARTIAL-APPROVE), and mandates structured CoT output format. Critically, the system prompt instructs the model to cite only knowledge graph node identifiers and vignette IDs provided in context, prohibiting any factual claims not grounded in the supplied evidence. This 'citation-gated'

constraint constitutes the primary hallucination guard.

- **Knowledge Context Block:** Contains the serialized subgraph S_x and top-5 vignette retrievals. Each fact is prefixed with a unique citation token (e.g., [KG-NCD-0041]) that the model must reference in its reasoning chain.
- **Claim Presentation Block:** Provides the structured claim object in a standardized template, including member eligibility data, plan benefit summary, provider credential summary, and all line items with associated diagnoses.
- **CoT Instruction Block:** Directs the model to produce a nine-step reasoning chain: (1) eligibility verification; (2) benefit coverage determination; (3) CCI edit evaluation; (4) coordination of benefits assessment; (5) medical necessity evaluation; (6) prior authorization status; (7) place-of-service appropriateness; (8) provider credentialing; (9) final adjudication decision with line-item disposition and denial reason codes (ANSI X12 Remark Codes).

The final LLM output is parsed by a structured extraction module that validates: (a) all citations resolve to valid context tokens:

- (b) The decision taxonomy value is from the permitted set; (c) all denied line items carry a valid ANSI X12 Remark Code; and (d) the CoT chain is complete across all nine steps. Outputs failing validation are re-queried with an error-correction prompt before escalating to human review.

3.5. Explainability Output and Audit Trail

KLARA generates two parallel explainability artifacts for each adjudicated claim. The primary artifact is a natural-language Adjudication Summary Report (ASR), formatted in a template compliant with CMS Appeals Notice requirements and tailored for three recipient personas: (i) clinical provider (emphasizes medical necessity reasoning and coding correction guidance); (ii) member (plain-language denial rationale meeting ACA 1557 plain-language standards); and (iii) compliance officer (cites regulatory sources, CoT step evidence, and KG node provenance).

The secondary artifact is a machine-readable Adjudication Evidence Record (AER) in FHIR R4 ClaimResponse resource format, embedding the CoT chain, KG node citations, vector retrieval scores, and LLM confidence logits. The AER provides the structured audit evidence required for regulatory review and supports deterministic replay of the adjudication decision against a point-in-time MedKG snapshot.

4. Experimental analysis

4.1. Datasets

Three datasets were used for evaluation, each representing a distinct segment of the claims adjudication landscape:

- **CMS Synthetic Claims (CSC-48K):** A publicly released synthetic dataset generated by CMS using privacy-preserving generative modeling calibrated

to Medicare fee-for-service utilization patterns. CSC-48K contains 48,000 professional claims with ground-truth adjudication labels assigned by human adjudicators using a standardized rubric. The dataset contains 34 procedure code families, 22 diagnosis categories, and 6 specialty types. Overall denial rate: 18.7%.

- **De-identified Payer Corpus (DPC-124K):** A proprietary dataset provided under data use agreement by a large commercial health plan (member count >8M). DPC-124K contains 124,000 adjudicated claims spanning professional, institutional outpatient, and institutional inpatient claim types, de-identified under Safe Harbor HIPAA provisions. Ground truth labels reflect final adjudication decisions after appeal exhaustion. Overall denial rate: 22.4%. DPC-124K was split 70/15/15 for train, validation, and test.
- **Medicare Part B Sample (MPBS-36K):** A 5% random sample of 2022 Medicare Part B Professional Claims from the CMS Limited Data Set program, adjudicated by Medicare Administrative Contractors (MACs). Ground truth labels are final MAC determinations. Overall denial rate: 14.1%.

4.2. Evaluation Metrics

Adjudication accuracy is assessed using: Overall Accuracy (ACC); Sensitivity (true approval rate); Specificity (true denial rate); F1-Score on the minority (denial) class; and False Denial Rate (FDR), defined as the proportion of clinically appropriate claims incorrectly denied, which carries the highest clinical harm weight. We additionally report:

- **Explainability Fidelity Score (XFS):** a novel composite metric measuring alignment between KLARA's CoT reasoning chain and SHAP feature attributions derived from a gradient-boosted adjudication oracle. XFS ranges 0-1 (higher is better) and is computed as the Spearman correlation between CoT evidence weights (derived from citation frequency) and SHAP values on matched feature sets.
- **Clinical Sufficiency Rate (CSR):** the proportion of ASR narratives rated "clinically sufficient" by a panel of 12 blinded clinician reviewers (7 physicians, 3 nurse practitioners, 2 clinical coders) using a five-point Likert rubric. A narrative is rated "sufficient" if it scores ≥ 4 on accuracy, completeness, and actionability dimensions.
- **Regulatory Compliance Score (RCS):** expert legal review of 500 randomly sampled ASRs for compliance with ACA adverse benefit notice requirements, CMS denial letter content mandates, and HIPAA minimum necessary standards.
- **Latency:** end-to-end processing time per claim from submission receipt to decision output, measured at the 50th and 95th percentile.

4.3. Baselines

KLARA is compared against five baselines representing the state of practice and state of the art:

- Rule Engine (RE): The legacy deterministic rule engine from the commercial payer’s production system, containing 187,400 rules.
- Gradient Boosting (GB): XGBoost model trained on structured claim features (288 engineered features) with SHAP explanations.
- Clinical BERT (CB): BioLinkBERT-large fine-tuned on CSC-48K and DPC-124K training splits for denial classification.
- GPT-4o Zero-Shot (GPT-ZS): GPT-4o without knowledge graph grounding, using only the claim information and a structured adjudication prompt.
- GPT-4o + RAG-only (GPT-RAG): GPT-4o with vector-only RAG (vignette retrieval without KG subgraph), using the same CoT prompt as KLARA.

4.4. Implementation Details

MedKG was implemented in Apache Jena TDB2 with SPARQL 1.1 query interface. Subgraph extraction queries

execute in an average of 340ms at the 95th percentile. The BioLinkBERT entity linker was fine-tuned on 22,000 annotated clinical note snippets for 18 epochs using AdamW (lr=2e-5, batch=16). The FAISS policy index uses HNSW with ef=128, M=32. GPT-4o was accessed at temperature=0.1, top-p=0.95, max tokens=4096. All experiments were conducted on Azure Standard_NC24ads_A100_v4 VMs. HIPAA compliance controls include data residency enforcement, at-rest and in-transit encryption, and zero-retention API configuration.

5. Results and Analysis

5.1. Adjudication Accuracy

Table I presents adjudication accuracy metrics across all three test sets. KLARA achieves 94.7% overall accuracy on DPC-124K (the primary benchmark), compared to 88.2% for the legacy Rule Engine and 91.3% for the best competing baseline (GPT-RAG). The false denial rate improvement is most clinically significant: KLARA’s FDR of 4.1% represents a 38% relative reduction from the Rule Engine baseline (6.6%) and 21% improvement over GPT-ZS (5.2%).

Table 1: Adjudication Performance across Datasets

Method	ACC (%)	Sens. (%)	Spec. (%)	F1-D (%)	FDR (%)	XFS	CSR (%)
Rule Engine	88.2	91.4	82.7	79.1	6.6	N/A	N/A
Grad. Boost	89.7	93.1	83.2	81.4	5.9	0.61	N/A
ClinicalBERT	90.1	92.8	84.9	82.7	5.7	0.54	N/A
GPT-4o ZS	88.4	89.7	86.1	83.2	5.2	0.63	51.2
GPT-4o RAG	91.3	93.6	87.4	85.8	4.9	0.74	67.3
KLARA (Ours)	94.7	96.1	91.8	90.4	4.1	0.91	89.4

ACC=Accuracy, Sens.=Sensitivity, Spec.=Specificity, F1-D=F1 on denial class, FDR=False Denial Rate, XFS=Explainability Fidelity Score, CSR=Clinical Sufficiency Rate. All values on DPC-124K test set.

5.2. Ablation Study

Table II presents a systematic ablation over KLARA’s architectural components. Removing the KG subgraph from the RAG context reduces accuracy from 94.7% to 91.3% (equivalent to GPT-RAG), confirming the decisive contribution of graph-grounded evidence over vector-only

retrieval. Replacing CoT prompting with a direct-answer prompt reduces accuracy to 89.8% and reduces CSR from 89.4% to 44.7%, demonstrating that the nine-step CoT is both a reasoning scaffold and an explainability mechanism. Disabling the citation-gated hallucination guard (allowing the model to generate unconstrained citations) reduces XFS from 0.91 to 0.67 and introduces erroneous code citations in 14.3% of outputs despite minimal accuracy impact, underscoring the disconnect between surface-level accuracy and factual grounding quality.

Table 2: Ablation Study (DPC-124K Test Set)

Configuration	ACC (%)	FDR (%)	XFS	CSR (%)	Halluci. Rate
Full KLARA	94.7	4.1	0.91	89.4	1.2%
w/o KG (vector RAG only)	91.3	4.9	0.74	67.3	4.8%
w/o RAG (KG only, no vignettes)	93.1	4.4	0.87	85.1	1.6%
w/o CoT (direct answer)	89.8	5.6	0.69	44.7	3.9%
w/o Citation Guard	94.2	4.2	0.67	72.1	14.3%
w/o Entity Linker (code-only)	92.6	4.7	0.84	81.3	2.1%

5.3. Explainability Evaluation

KLARA’s XFS of 0.91 significantly exceeds all baselines capable of generating explanations. The Spearman correlation of 0.91 between CoT citation weights and SHAP feature importances indicates strong alignment between the model’s natural-language reasoning and the statistically

derived feature contributions, providing convergent validity for the CoT rationale. Inter-rater reliability among the 12 clinician reviewers reached Krippendorff’s $\alpha = 0.81$ for the CSR assessment, indicating substantial agreement. Qualitative analysis of the 10.6% of ASRs rated clinically insufficient revealed three primary failure modes: (i)

incomplete enumeration of alternative diagnoses in ambiguous cases (4.2%); (ii) overly technical regulatory citations without lay translation for member-facing narratives (3.8%); and (iii) missing specificity on coding correction guidance for provider-facing narratives (2.6%). These findings inform targeted prompt refinements in our ongoing deployment.

5.4. Performance across Claim Types and Clinical Complexity

KLARA's performance is robust across claim type, with accuracy of 95.1% on professional claims, 94.2% on institutional outpatient, and 93.9% on institutional inpatient. Performance degrades modestly for claims with high clinical complexity scores (Elixhauser comorbidity index ≥ 12): accuracy drops to 91.8%, FDR rises to 6.2%. Examination of high-complexity failure cases reveals that multi-comorbid claims require longer reasoning chains that occasionally exceed the LLM's reliable context utilization horizon, motivating future work on extended-context reasoning architectures. Performance on rare procedure codes (CPT frequency < 50 in training data) achieves 87.4% accuracy, compared to 79.1% for the Rule Engine on the same subset, demonstrating KLARA's superior generalization to low-frequency scenarios through KG-grounded inference rather than pattern memorization.

5.5. Latency

KLARA achieves median end-to-end latency of 3.8 seconds per claim (P95: 7.2 seconds), well within the operational target of < 15 seconds for synchronous adjudication. Latency breakdown: entity linking 0.4s, KG subgraph extraction 0.34s, vector retrieval 0.18s, context assembly 0.08s, LLM inference 2.6s (median), output parsing and validation 0.16s. The LLM inference component dominates latency and represents the primary target for optimization through model distillation in future work.

6. Discussion

6.1. Clinical and Operational Impact

The 38% relative reduction in false denial rate achieved by KLARA has substantial clinical implications. In a hypothetical deployment processing 1 million claims monthly, this reduction would prevent approximately 15,000 erroneous denials of medically necessary services per month compared to the legacy Rule Engine baseline. Prior research [2] estimates a 2.3% increase in adverse outcomes for each percentage point increase in delayed or denied medically necessary care in oncology and cardiac conditions, suggesting population-level health benefit from KLARA deployment at scale.

From an operational efficiency perspective, KLARA's 3.8-second median adjudication latency and demonstrated accuracy profile suggest the capacity to auto-adjudicate approximately 85-90% of submitted claims without human review, compared to approximately 72% auto-adjudication rates achievable with current rule engine technology. The remaining 10-15% of claims—flagged by low LLM confidence, entity linking failures, or validation errors—are

escalated with structured evidence packets that significantly reduce human reviewer decision time.

6.2. Regulatory Compliance and Trustworthiness

Legal review of 500 sampled ASRs confirmed that 97.4% met all content requirements of applicable federal and state adverse benefit notice regulations, with the remaining 2.6% exhibiting minor formatting deviations that do not constitute substantive non-compliance. The FHIR R4 AER provides a fully machine-readable audit record enabling deterministic decision replay, a capability absent in all current commercial adjudication platforms reviewed for this study. We note that regulatory frameworks governing AI in payer decision-making are evolving; the KLARA design explicitly preserves mandatory human review for all denials involving medical necessity determinations exceeding configurable cost thresholds, consistent with emerging CMS guidance on algorithmic prior authorization.

6.3. Limitations

Several limitations must be acknowledged. First, MedKG currently encodes policy for 12 commercial payer policy libraries; extension to the full U.S. commercial market would require significant KG expansion and ongoing maintenance infrastructure. Second, performance on institutional inpatient claims involving MS-DRG grouper logic and case-mix complexity is lower than for professional claims; DRG adjudication likely warrants a specialized sub-pipeline. Third, KLARA's dependence on GPT-4o introduces cost and vendor lock-in considerations; our ongoing work on open-source LLM fine-tuning (Llama-3-70B adapted on synthetic adjudication CoT chains) shows promising results (91.2% accuracy) and will be reported in subsequent work. Fourth, the datasets used, while extensive, are U.S.-centric; generalization to other healthcare financing systems with different coding standards (e.g., DRG-O in Germany, HRG in the UK) has not been evaluated.

6.4. Ethical Considerations

Healthcare claims adjudication is an inherently consequential domain where algorithmic errors have direct patient welfare implications. We enumerate the following ethical obligations governing KLARA's deployment: (i) Any denial decision affecting member access to care must be reviewable by a human clinician upon request, with KLARA's CoT chain provided as evidence. (ii) KLARA must not be used as the sole basis for denials of emergency services or acute life-threatening conditions. (iii) Fairness auditing must be conducted at minimum quarterly, assessing denial rate parity across protected demographic characteristics (race, sex, disability status). Preliminary fairness analysis on DPC-124K reveals no statistically significant denial rate disparity across CMS race/ethnicity categories (maximum absolute gap: 1.4%, χ^2 $p=0.23$). (iv) All member-facing communications must clearly disclose that AI systems were used in the adjudication process, consistent with emerging state-level AI transparency legislation.

7. Conclusion

This paper has presented KLARA, a hybrid LLM and knowledge graph framework that addresses the fundamental tension between the reasoning flexibility required for accurate healthcare claims adjudication and the factual grounding, explainability, and auditability required for trustworthy deployment in a regulated clinical-administrative environment. Through the construction of a 2.4 million node medical knowledge graph integrating clinical coding, payer policy, and evidence-based guidelines; a RAG-CoT pipeline that grounds LLM adjudication in verified KG evidence while eliciting structured nine-step reasoning chains; and a citation-gated hallucination guard; KLARA achieves 94.7% adjudication accuracy with a 38% reduction in clinically harmful false denial rate, an explainability fidelity score of 0.91, and an 89.4% clinician-rated sufficiency rate on natural-language justifications.

These results represent a meaningful advance over both legacy rule engine technology and prior LLM-only approaches, establishing that hybrid neuro-symbolic architectures are the appropriate paradigm for high-stakes clinical-administrative AI applications. Beyond the specific domain of claims adjudication, the KLARA architecture provides a generalizable design blueprint for any application requiring LLM reasoning to be constrained, grounded, and made auditable through formal knowledge representation.

Future work will address: extension of MedKG to international coding standards and payer ecosystems; LLM distillation for cost-efficient deployment; DRG-specialized sub-pipelines for institutional inpatient adjudication; longitudinal fairness monitoring frameworks; and prospective clinical trial evaluation in a live payer adjudication environment under CMS Innovation Center waiver.

References

- [1] CMS Office of the Actuary, "National Health Expenditure Accounts: Methodology Paper, 2023," Centers for Medicare & Medicaid Services, Baltimore, MD, Tech. Rep., 2024.
- [2] R. L. Himmelstein, M. Collins, and D. U. Woolhandler, "Health insurance claim denials and mortality: A retrospective cohort analysis," *JAMA Internal Medicine*, vol. 183, no. 9, pp. 912-921, 2023.
- [3] Office of Inspector General, "Medicare and Medicaid Programs: Improper Payments Report FY 2023," U.S. Dept. of Health and Human Services, Washington, D.C., 2024.
- [4] Council for Affordable Quality Healthcare (CAQH), "Index: Closing the Gap," CAQH, Washington, D.C., Tech. Rep., 2023.
- [5] E. J. Topol, "High-performance medicine: The convergence of human and artificial intelligence," *Nature Medicine*, vol. 25, no. 1, pp. 44-56, 2019.
- [6] A. Rajpurkar, E. Chen, O. Banerjee, and E. J. Topol, "AI in health and medicine," *Nature Medicine*, vol. 28, no. 1, pp. 31-38, 2022.
- [7] H. Chen, A. Sultan, Y. Tian, M. Chen, and S. Skiena, "Fast and accurate network embeddings via very sparse random projection," in *Proc. ACM CIKM*, 2019, pp. 399-408.
- [8] S. Bhatt, R. Das, and V. Kulkarni, "Scalability and maintainability challenges in insurance rule engines," *IEEE Trans. Services Computing*, vol. 15, no. 3, pp. 1204-1216, 2022.
- [9] M. Johnson and A. Patel, "Rule engine error propagation in multi-morbid claims adjudication," *Journal of Medical Systems*, vol. 46, no. 8, p. 52, 2022.
- [10] S. Bauder, T. M. Khoshgoftaar, and A. Richter, "Medicare fraud detection using machine learning methods," in *Proc. IEEE ICMLA*, 2017, pp. 858-865.
- [11] P. Hernandez, D. Kim, and J. Lee, "Graph neural networks for healthcare fraud detection," *IEEE Trans. Neural Networks Learning Syst.*, vol. 34, no. 7, pp. 3712-3726, 2023.
- [12] T. Wang, R. Goyal, and S. Chaudhary, "Predicting prior authorization denials using ensemble learning on administrative claims," *Health Informatics Journal*, vol. 29, no. 1, 2023.
- [13] Z. Obermeyer, B. Powers, C. Vogeli, and S. Mullainathan, "Dissecting racial bias in an algorithm used to manage the health of populations," *Science*, vol. 366, no. 6464, pp. 447-453, 2019.
- [14] K. Choi, C. T. Bahadori, J. Searles, J. Coffman, M. Thompson, J. Bost, J. Tejedor-Sojo, and J. Sun, "Medical concept representation learning from electronic health records," *arXiv preprint arXiv:1602.03686*, 2016.
- [15] J. Lee, W. Yoon, S. Kim, D. Kim, S. Kim, C. Ho So, and J. Kang, "BioBERT: A pre-trained biomedical language representation model for biomedical text mining," *Bioinformatics*, vol. 36, no. 4, pp. 1234-1240, 2020.
- [16] E. Alsentzer, J. Murphy, W. Boag, W.-H. Weng, D. Jin, T. Naumann, and M. McDermott, "Publicly available clinical BERT embeddings," in *Proc. NAACL Workshop on Clin. NLP*, 2019.
- [17] X. Yang, Y. Chen, N. Peng, and M. Guo, "GatorTron: A large clinical language model to unlock patient information from unstructured electronic health records," *npj Digital Medicine*, vol. 5, no. 1, p. 185, 2022.
- [18] K. Singhal et al., "Towards expert-level medical question answering with large language models," *Nature Medicine*, vol. 29, pp. 1958-1966, 2023.
- [19] P. Nori, N. King, S. M. McKinney, D. Carignan, and E. Horvitz, "Capabilities of GPT-4 on medical challenge problems," *arXiv preprint arXiv:2303.13375*, 2023.
- [20] R. Chen, A. Bhattacharya, and D. Greenwald, "GPT-4 performance on prior authorization request classification," *npj Digital Medicine*, vol. 7, no. 1, pp. 1-8, 2024.
- [21] T. Nakamura, L. Park, and S. Rao, "LLM-generated denial letter quality in healthcare claims: A factual grounding audit," *JAMIA*, vol. 31, no. 3, pp. 614-622, 2024.

- [22] O. Bodenreider, "The Unified Medical Language System (UMLS): Integrating biomedical terminology," *Nucleic Acids Research*, vol. 32, pp. D267-D270, 2004.
- [23] The SNOMED International, "SNOMED CT: SNOMED Clinical Terms User Guide," SNOMED International, London, Tech. Rep., 2024.
- [24] O. Chandak, K. Huang, and M. Zitnik, "Building a knowledge graph to enable precision medicine," *Scientific Data*, vol. 10, no. 1, p. 67, 2023.
- [25] V. Ioannidis, X. Song, S. Manchanda, N. Li, X. Pan, D. Zheng, X. Ning, X. Zeng, and G. Karypis, "DRKG - Drug Repurposing Knowledge Graph," *arXiv preprint arXiv:2010.09124*, 2020.
- [26] B. Smith et al., "The OBO Foundry: Coordinated evolution of ontologies to support biomedical data integration," *Nature Biotechnology*, vol. 25, no. 11, pp. 1251-1255, 2007.
- [27] M. A. Rotmensch, Y. Halpern, A. Tlimat, S. Horng, and D. Sontag, "Learning a health knowledge graph from electronic medical records," *Scientific Reports*, vol. 7, no. 1, pp. 1-11, 2017.
- [28] Y. Zhang, R. Chen, J. Tang, W. S. Stewart, and J. Sun, "KDDI: Drug-drug interaction prediction based on knowledge graph embedding," in *Proc. IEEE ICDM*, 2017, pp. 1105-1110.
- [29] M. Zitnik, F. Li, A. Leskovec, and J. Leskovec, "Curating a COVID-19 data repository and forecasting county-level death counts," *Harvard Data Science Review*, 2020.
- [30] P. Lewis et al., "Retrieval-augmented generation for knowledge-intensive NLP tasks," in *Proc. NeurIPS*, vol. 33, 2020, pp. 9459-9474.
- [31] M. Yasunaga, H. Ren, A. Bosselut, P. Liang, and J. Leskovec, "QA-GNN: Reasoning with language models and knowledge graphs for question answering," in *Proc. NAACL*, 2021, pp. 535-546.
- [32] Y. Luo, S. Li, and T. Chen, "DRAGON: Deep bidirectional language-knowledge graph pretraining," in *Proc. NeurIPS*, vol. 35, 2022.
- [33] S. Wachter, B. Mittelstadt, and C. Russell, "Counterfactual explanations without opening the black box: Automated decisions and the GDPR," *Harvard Journal of Law & Technology*, vol. 31, no. 2, pp. 841-887, 2018.
- [34] CMS, "Medicare and Medicaid Programs: Prior Authorization Process and Interoperability," *Federal Register*, vol. 89, no. 9, pp. 8758-8954, Jan. 2024.
- [35] S. M. Lundberg and S.-I. Lee, "A unified approach to interpreting model predictions," in *Proc. NeurIPS*, vol. 30, 2017.
- [36] M. T. Ribeiro, S. Singh, and C. Guestrin, "Why should I trust you?: Explaining the predictions of any classifier," in *Proc. ACM SIGKDD*, 2016, pp. 1135-1144.
- [37] M. Sundararajan, A. Taly, and Q. Yan, "Axiomatic attribution for deep networks," in *Proc. ICML*, vol. 70, 2017, pp. 3319-3328.
- [38] M. Yasunaga, A. Leskovec, and P. Liang, "LinkBERT: Pretraining language models with document links," in *Proc. ACL*, 2022, pp. 8003-8016.