

AUTOMATED BIM LOD INFERENCING: A HYBRID APPROACH USING KNOWLEDGE GRAPHS AND MULTIMODAL LARGE LANGUAGE MODELS

Taeyoon Kim¹, Ghang Lee¹

¹Yonsei University, Seoul, Republic of Korea

Abstract

Automated Level of Development (LOD) assessment in Building Information Modeling (BIM) is hindered by semantic ambiguity in specifications and geometric complexity of objects. Existing rule-based and vision-based methods lack contextual reasoning or struggle with fragmented, occluded geometry. This study proposes a hybrid framework combining a three-layered knowledge graph—built top-down via LLM requirement extraction and bottom-up via visual population—with a geometric pipeline reassembling meshes and exposing occluded components. A multimodal LLM cross-references graph-derived criteria with multi-view geometry, producing explainable assessments. A four-condition ablation study achieves 83.47% accuracy while reducing false negatives, confirming feasibility for explainable BIM quality assurance.

Introduction

Research Background

As digital transformation accelerates within the construction industry, Building Information Modeling (BIM) has established itself as a core database for managing information throughout the entire project lifecycle, moving beyond simple 3D geometric modeling. To maximize the utility of BIM data, it is essential to verify the accuracy of the information and ensure compliance with the Level of Development (LOD) standards required for each project phase. However, existing LOD verification methods heavily rely on manual inspection by experts, which incurs significant time and costs. Furthermore, conventional rule-based automation tools lack the flexibility to handle various geometric variations and complex exceptions. While recent studies have attempted to leverage Artificial Intelligence (AI) to address these issues, two fundamental barriers remain: the "semantic gap" and "geometric complexity."

Problem Statement

The first challenge is the semantic gap between human language and machine understanding. LOD specifications are written in natural language, possessing linguistic characteristics that are difficult for machines to interpret intuitively. For instance, umbrella terms such as "All Reinforcement" do not map one-to-one with specific

objects (e.g., main bars, stirrups) within the model. Additionally, context-dependent terms like "Connection" exhibit polysemy, where the target object for inspection changes depending on whether the system is Steel or Precast Concrete. Simple text-matching algorithms may utilize static dictionaries, but they fail to dynamically infer these hierarchical and contextual relationships, preventing the scalable application of automated verification.

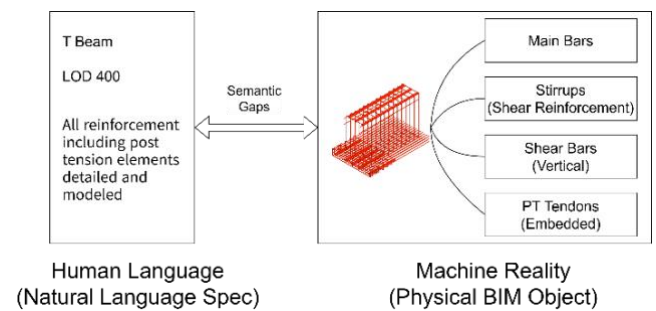


Figure 1 The semantic gap in LOD verification.

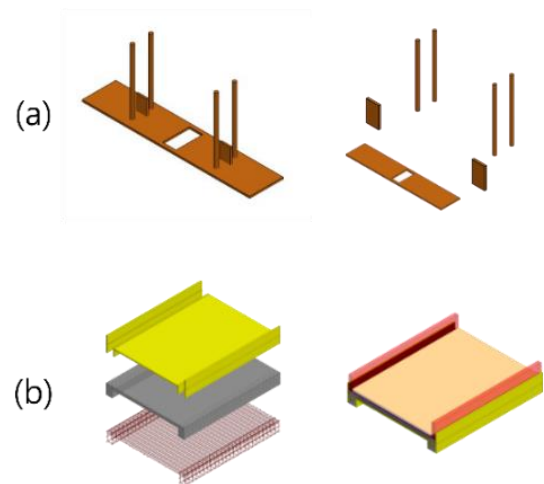


Figure 2 Visual illustration of geometric complexity in BIM objects. (a) Fragmentation: A single embedded component is geometrically constructed from multiple disconnected meshes.

(b) Occlusion: A baseplate connection hides internal reinforcement behind opaque surfaces, impeding visual identification.

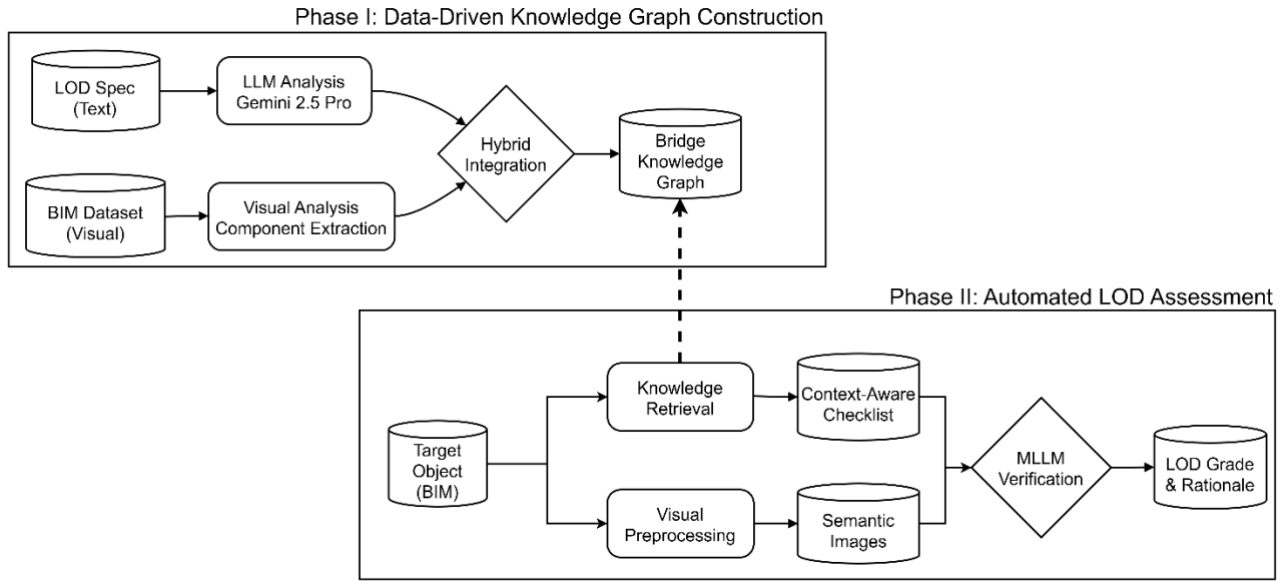


Figure 3 The architecture of the proposed dual-phase LOD assessment framework

The second challenge is geometric complexity. Although a BIM object may logically represent a single element (e.g., a Precast Concrete T-beam), its sub-components often exhibit geometric fragmentation. For instance, a single embedded steel part is frequently composed of multiple disconnected meshes—where the base plate and studs are modeled as separate geometric entities—lacking a unified topological structure (fragmentation). Conversely, objects often contain critical internal components, such as reinforcement within a baseplate connection, that are visually obscured by external surfaces (occlusion). If such objects are simply captured as raw images, the AI model is likely to misinterpret the fragmented meshes as separate unrelated objects, while failing to recognize the hidden internal components.

Research Objective and Organization

The objective of this research is to establish an intelligent quality assessment system that integrates domain knowledge with visual reasoning. Unlike black-box AI approaches, this study proposes a knowledge graph-based framework to ensure explainable assessments. Specifically, the system aims to (1) resolve linguistic ambiguities in LOD standards by constructing a logical reasoning engine, and (2) overcome visual occlusion and fragmentation through semantic geometric preprocessing. By bridging these semantic and geometric gaps, the framework enables MLLMs to perform rigorous and traceable LOD verification.

Related Works

This section reviews existing literature on BIM object classification and quality assessment, categorizing them into (1) geometry and metadata-based approaches, (2) deep Learning-based visual classification, and (3) Recent MLLM applications. It highlights the limitations of each approach and defines the differentiation of the proposed research.

Geometry and Metadata-based Approaches

Early studies focused on utilizing geometric features or textual metadata of BIM objects for classification. (Wu and Zhang, 2019) proposed a hybrid algorithm combining geometric features with keyword matching to enhance BIM interoperability. However, this method relies heavily on pre-defined rules, lacking flexibility for irregular objects that deviate from standard specifications. Similarly, (Abualdenien and Borrmann, 2022) utilized ensemble learning to classify the Level of Geometry (LOG) based on geometric complexity (e.g., area, volume). While effective for assessing geometric precision, this approach fails to semantically identify the presence of specific components (e.g., bolts, nuts), which is the core requirement for actual LOD assessment.

Vision-based Deep Learning and Data Dependency

To overcome the limitations of metadata reliance, research on converting 3D BIM objects into 2D images for classification using Convolutional Neural Networks (CNN) has gained traction. (Yang et al., 2026) improved classification performance for architectural and MEP objects using multi-view images. However, these Supervised Learning approaches entail a significant critical limitation: they require the costly pre-construction of massive, labeled datasets containing thousands of images.

To address performance issues, (Yu et al., 2023) and (Utkucu et al., 2024) proposed Ensemble Deep Learning techniques and Random Forest integration to enhance accuracy and compatibility with classification systems like Uniclass. Despite these improvements, they still fail to overcome the fundamental "Data Hunger" problem; whenever a new LOD standard or a new object type is introduced, new data must be collected, and the model must be re-trained. Furthermore, relying solely on exterior visual features prevents these models from identifying

occluded components—such as rebar inside concrete—which are essential for detailed LOD verification.

Multimodal LLMs and Ontology Approaches in Generative Interpretation

Recently, studies utilizing MLLMs have emerged to achieve Zero-shot classification and reduce data construction costs. (Wang, X. et al., 2025) demonstrated the use of models like GPT-4V to detect surface cracks on steel bridges, and (Zeng et al., 2024) estimated building age solely from facade images. While these studies show the potential of MLLMs, they are limited to surface-level image analysis. They have not attempted to integrate domain knowledge (knowledge graph) to evaluate objects with complex hierarchies and contexts, such as BIM elements. Without structured knowledge, general-purpose MLLMs are prone to hallucinations and fail to grasp the specific engineering nuances required for LOD assessment.

Recent advancements in Natural Language Processing (NLP) and ontology have significantly improved the structuring of unstructured construction data, primarily enhancing information extraction and retrieval tasks. For instance, (Nahri et al., 2025) utilized transformer-based models to extract explicit requirements from technical specifications, while (Yin et al., 2023) developed an ontology-aided approach to translate natural language queries for retrieving information from existing BIM models. Similarly, (F.h et al., 2025) formalized the ISO 19650 standard into the BIM-OIM ontology to support process compliance. While these studies effectively parse explicit text or validate established processes, the domain of generative modeling presents an opportunity for further exploration regarding the "Semantic gap." Current methods typically focus on direct mapping, and there is potential to extend these approaches to include the inferential reasoning required to translate abstract "Human Standards" (e.g., vague LOD descriptions) into concrete "Deep Standards" (e.g., specific component hierarchies). This study aims to build upon these foundations by introducing a semantic bridge to infer implied domain concepts—such as deducing required Connection Assemblies from a text that mentions anchor rods—thereby facilitating the automated decision-making process in BIM object creation.

Methodology

Framework Overview

This study proposes a comprehensive framework designed to bridge the semantic gap between human requirements and machine interpretation for automated LOD verification. The proposed system is structurally divided into two primary phases: (1) Phase I: Data-Driven knowledge graph Construction and (2) Phase II: Automated LOD Assessment. In the first phase, a Bridge knowledge graph (KG) is constructed as a foundational reasoning engine. Instead of relying on static rule-based definitions, this knowledge base is developed through a hybrid data-driven process where natural language requirements are parsed by LLM, specifically Gemini 2.5

Pro, while physical component vocabularies are derived from the visual analysis of high-fidelity BIM data. In the second phase, the system executes the assessment by cross-referencing this pre-validated knowledge with the geometric data of the target BIM object. This dual-phase approach ensures both the scalability of knowledge acquisition and the consistency of real-time verification.

Phase I: Data-Driven knowledge graph Construction

Three-Layer Graph Architecture

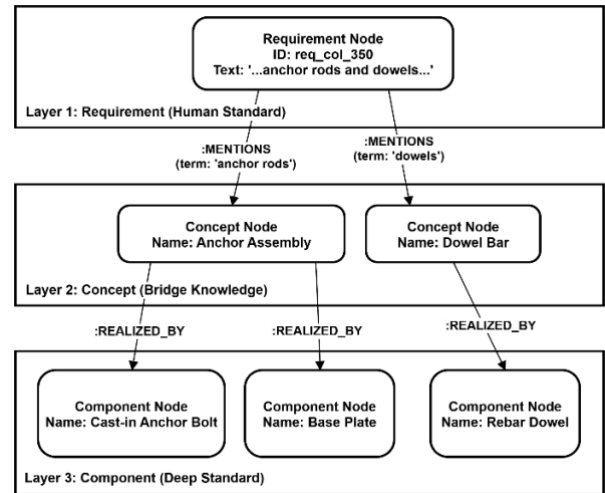


Figure 4 The three-layered architecture of the Bridge knowledge graph.

A critical limitation of existing LOD verification methods is the inability to rigorously map abstract natural language requirements to specific 3D objects. In this study, we distinguish between two levels of specification: "Human Standards," defined as the natural-language requirements stated in LOD specifications (e.g., "Include anchor rods and dowels"), which are intended for human interpretation and often underspecified for automated inspection. In contrast, "Deep Standards," denote a machine-interpretable representation of these requirements, expressed as structured inspection criteria that explicitly define each required sub-component in terms of its identity and expected geometric form (e.g., "Cast-in Anchor Bolt: cylindrical rod embedded in concrete"). The term "Deep" emphasizes that these standards capture latent, operational constraints implicit in human-readable specifications but not explicitly stated. The core objective of the proposed knowledge graph is to bridge this gap by systematically translating Human Standards into Deep Standards, enabling reliable reasoning and verification by downstream visual and geometric analysis modules.

To achieve this, we designed a three-layer graph architecture, as illustrated in Figure 4. The top layer, Requirement Nodes, encapsulates the raw textual requirements extracted from standard specifications. For example, as shown in the diagram, a requirement for LOD 350 Columns might state: "Include anchor rods and dowels." The bottom layer, Component Nodes, consists of concrete physical elements existing in the modeling

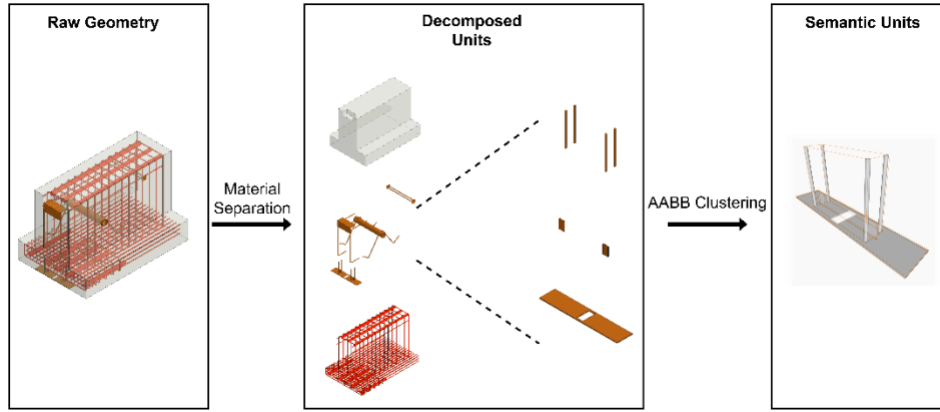


Figure 5 Transformation from raw geometry to semantic units. Visualization of the preprocessing steps: Raw input, Decomposed units via material separation, and Re-assembled semantic units via AABB clustering.

environment, such as "cast-in anchor bolt," "base plate," and "rebar dowel." Connecting these two is the Concept Node Layer (Bridge Layer). This intermediate layer acts as a semantic hub that abstracts ambiguous terms into system-specific engineering concepts. In Figure 4, the term "anchor rods" from the requirement is mapped to the concept "anchor assembly" (via the "mentions" relationship), which is then physically realized by components like "Cast-in Anchor Bolt" and "Base Plate" (via the "realized_by" relationship). This structure effectively resolves linguistic ambiguity by integrating diverse expressions into unified technical definitions.

It should be noted that the semantic mapping between layers in the current knowledge graph was manually curated. While LLMs assisted in extracting requirements and suggesting candidate concepts, the final alignment was performed through manual validation. Consequently, when a new LOD standard is introduced, the Requirement layer must be reconstructed, whereas the Concept and Component layers—representing reusable engineering vocabulary—can be largely retained and extended incrementally.

Hybrid Learning Pipeline

To populate this three-layer structure, we employed a Hybrid Learning Pipeline that synthesizes both textual and visual data. First, adopting a top-down approach, we utilized GPT-5.1 to analyze unstructured PDF standard documents. The model extracts core requirements and infers implied engineering concepts (implied concepts) to construct the Requirement and Concept layers. Simultaneously, a Bottom-Up approach was applied using Visual Analysis techniques on a dataset of high-fidelity (LOD 400) BIM models. By scanning these models, the system extracted a "Visual Vocabulary" of actual existing components to populate the Component Layer. Finally, through a Semantic Alignment process, these two streams were mapped to complete a validated knowledge graph that eliminates the discrepancy between "what is required" (Human Standards) and "what is modeled" (Deep Standards).

Phase II: Context-Aware & Explainable Assessment

Context-Aware Knowledge Retrieval

During the verification phase, the system utilizes the pre-constructed knowledge graph as a deterministic reasoning engine. When a user selects a BIM object, the system queries the graph using the object's classification code (e.g., OmniClass). The query engine performs Context-Aware Retrieval, dynamically interpreting ambiguous terms based on the object's system type—for instance, mapping the term "Connection" to "Anchor Bolts" for a steel column, while mapping it to "Shear Tabs" for a beam. This process instantly generates a precise, object-specific checklist without redundant inference computations.

Semantic Geometric Preprocessing

In typically complex BIM environments, internal components are often obscured by finishing materials (Occlusion) or exist as fragmented meshes (Fragmentation). To address this, the visual preprocessing module transforms raw geometry into machine-readable "Semantic Units," as illustrated in Figure 5.

The process begins with Raw Geometry, where internal reinforcement and embed plates are hidden within the concrete mass (Left). The module first executes Material Separation, decomposing the object into distinct layers—separating opaque concrete from internal steel elements—to generate Decomposed Units (Middle). This step effectively creates "X-ray Views" that reveal previously occluded structures.

Subsequently, the system employs Axis-Aligned Bounding Box (AABB) Clustering to reassemble these scattered parts. By analyzing spatial proximity, fragmented meshes (e.g., separate bolts and plates) are merged into unified Semantic Units (Right). This final representation provides the MLLM with a clear, obstacle-free view of the target components, ensuring accurate identification.

Evidence-Based MLLM Verification

Finally, Google's Gemini 2.5 Pro is employed as the central Multimodal LLM (MLLM) verification engine,

evaluating LOD compliance by cross-referencing the logical checklist derived from the knowledge graph with the preprocessed visual evidence. To prevent logical circularity—where the same model validates its own reasoning—this framework utilizes a cross-model verification strategy: OpenAI’s GPT-5.1 was employed during the initial phase to generate the logical rationales and Knowledge Graph, while Google’s Gemini 2.5 Pro is used independently for the final inspection phase. Unlike traditional black-box AI models that output opaque probability scores, the proposed system generates traceable rationales for its decisions. For instance, instead of merely labeling an object as "LOD 350," the system explicitly validates whether specific criteria—such as the visual presence of "Lifting Anchor Assemblies"—are met.

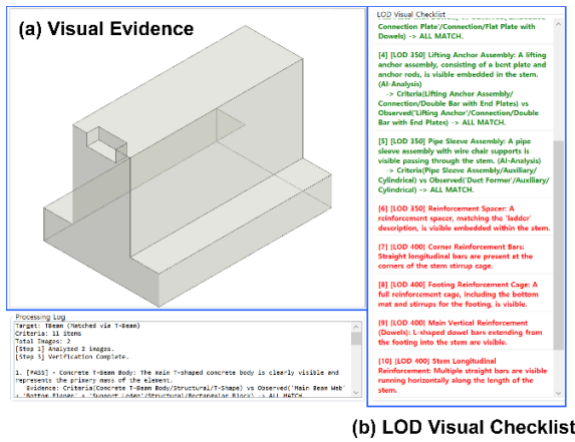


Figure 6 The integrated verification dashboard within Autodesk Revit.

To ensure seamless integration into the engineer's workflow, this verification logic is embedded directly within the BIM authoring environment via the Revit API, as illustrated in Figure 6. The system provides a real-time LOD Visual Checklist panel that displays the itemized inspection results derived from the knowledge graph. Compliant items are intuitively highlighted in green with an "ALL MATCH" tag, allowing users to instantly verify the satisfaction of specific requirements. This interactive feedback mechanism transforms the framework from a theoretical model into a practical decision-support tool, offering granular transparency beyond simple binary classifications.

Experimental Design and Validation Strategy

This section details the experimental framework established to assess the feasibility and performance of the proposed MLLM based LOD evaluation system. The composition of the dataset and the comparative scenarios designed to verify the impact of knowledge graph integration are described below.

Experimental Setup and Dataset

Dataset Composition

To evaluate the system's versatility across diverse construction disciplines, this study utilized the Level of Geometry (LOG) dataset (Abualdenien & Borrmann., 2022), a widely recognized benchmark in LOD research. The dataset encompasses a broad range of BIM families from structural, architectural, and MEP domains.

For the experimental scope, 31 representative families were selected, covering four LOD levels (LOD 200, 300, 350, and 400). This selection resulted in a total of 124 verification cases. These objects were specifically chosen to test the system's robustness against geometric complexities, including intricate internal components such as reinforcement cages, anchor bolts, and embed plates, which are rigorously modeled to comply with high-LOD standards.

Comparative Scenarios

To quantitatively measure the contribution of the proposed methodology, we conducted a comparative analysis between two distinct scenarios:

Baseline (Vanilla MLLM): This scenario represents the conventional verification method where raw text requirements from standard LOD Specifications are directly input into the MLLM without semantic processing. Under this baseline condition, an average of 4 criteria were applied per level, totaling 496 verification items. These criteria relied primarily on text-based descriptions, limiting the verification to abstract specifications rather than quantitative geometric analysis.

Proposed Method (Proposed KG+MLLM Approach): This scenario utilizes the Bridge knowledge graph constructed in this study. It applies the "deep standards checklist"—derived from the hybrid ontology learning pipeline—along with semantic visual preprocessing. This approach automatically generated 253 Concept Nodes and 626 Detailed Features. Consequently, the proposed method provides approximately 5.1 times more data points than the baseline, significantly enhancing the granularity of verification by grounding the process in actual model geometry.

Implementation Details

Model Implementation Details All experiments were conducted using the Gemini 2.5 Pro model via the API. This model was selected for its superior multimodal reasoning capabilities and extended context window, which are essential for processing the complex interplay between visual inputs (semantic units) and textual logic (knowledge graph queries).

Evaluation Metrics and Ground Truth

The quantitative evaluation employed confusion matrix metrics (TP, TN, FP, FN), determined based on the Cumulative Nature of LOD. This principle dictates that geometric features present in a lower LOD must persist in higher LOD levels.

To ensure reliability, the Ground Truth (GT) for comparison was established through a hybrid validation protocol. This approach utilized GPT-5.1 for initial automated verification, followed by rigorous manual inspection by researchers to finalize the dataset. The specific metrics are defined as follows:

True Positive (TP): System correctly classifies a Valid Model (e.g., Model 400) as Pass.

True Negative (TN): System correctly classifies an Invalid Model (e.g., Model 300 claiming to be 400) as Fail.

False Positive (FP): System incorrectly classifies an Invalid Model as Pass.

False Negative (FN): System incorrectly classifies a Valid Model as Fail.

Results and Discussion

Calibration: Threshold Sensitivity Analysis

Before performing the comparative analysis, a calibration experiment was conducted to determine the optimal “Decision Boundary” for the MLLM-based verification. Unlike rule-based systems, vision-based AI verification involves semantic ambiguity due to occlusion or resolution limits. Therefore, selecting an appropriate pass rate threshold is critical to balance sensitivity (Recall) and specificity (Precision).

Table 1: Sensitivity Analysis of Feature Emergence LOD Thresholds

Threshold	Self-Check (TP Rate)	Rejection (TN Rate)	Balance Score (TP+TN)
20%	99.2%	64.1%	1.63
40%	94.3%	85.3%	1.80 (Max)
60%	85.4%	93.5%	1.79
80%	65.9%	98.4%	1.64
100%	50.4%	98.9%	1.49

Table 2: Performance Comparison between Baseline and Proposed Framework

Metric	Vanilla MLLM	KG Only	Cluster IMG Only	KG+Cluster IMG
Accuracy	51.01%	67.34%	48.19%	83.47%
F1 Score	35.54%	66.39%	29.97%	86.15%
Precision	100.00%	93.02%	96.49%	90.43%
Recall	21.61%	51.61%	17.74%	82.26%
True Positive (TP)	67	160	55	255
True Negative (TN)	186	174	184	159
False Positive (FP)	0	12	2	27
False Negative (FN)	243	150	255	55

Table 1 presents the system’s performance across pass rate thresholds ranging from 20% to 100%. The “Self-Check” score indicates the probability of a valid model passing its own standard, while the “Rejection” score indicates the probability of a lower-LOD model failing a higher standard.

The analysis reveals that the 40% threshold serves as the optimal operating point, maximizing the Balance Score (1.80). At this threshold, the system correctly identifies 94.3% of valid models while successfully rejecting 85.3% of invalid claims. Conversely, stricter thresholds (80-100%) caused a sharp decline in Recall (down to 50.4%), confirming that requiring a 100% feature match is impractical due to visual occlusion in 2D renders.

Consequently, this study adopts 40% as the standard threshold for the final comparative analysis.

Comparative Performance Analysis

Table 2 presents the ablation study results under the calibrated 40% threshold across four conditions: Vanilla MLLM (baseline), KG Only, Cluster IMG Only, and KG+Cluster IMG (proposed).

The Vanilla baseline exhibited a Recall of only 21.61% with 243 False Negatives, confirming that unstructured LOD descriptions cause the MLLM to default to negative predictions. Its Precision of 100% (zero FPs) reveals an extreme conservative bias rather than genuine accuracy.

Introducing semantic decomposition alone (KG Only) raised Recall to 51.61% by breaking abstract criteria into ~5.1 concrete inspection items, but plateaued at 67.34% accuracy due to insufficient visual resolution from a single overview image. Conversely, providing cluster images without structured criteria (Cluster IMG Only) yielded the worst performance (Recall: 17.74%), even below the baseline—demonstrating that detailed visual inputs without a semantic framework for interpretation lead to near-systematic rejection.

The proposed KG+Cluster IMG framework achieved the highest performance across all metrics (Accuracy: 83.47%, F1: 86.15%, Recall: 82.26%), confirming a complementary synergy: the Knowledge Graph provides what to inspect, while cluster images provide where to

inspect. Neither alone is sufficient, but their combination boosts Recall by 30.65 percentage points over KG Only.

The ablation results further reveal that FPs originate primarily from semantic decomposition rather than visual input: the KG-Only case produced 12 FPs, compared with only 2 for the Cluster-IMG-Only case. The combined framework yields 27 FPs, representing a favorable trade-off, as it recovers 188 TPs at a 1:7 TP-TP ratio. This balance is operationally desirable in engineering quality assurance, where undetected deficiencies (FN) pose substantially greater risk than over-acceptance (FP).

Overall, the F1 Score surged from 35.54% to 86.15%, signifying a qualitative leap from unusable baseline performance to engineering-deployable reliability. The ablation confirms that automated LOD verification requires both structured semantic criteria and spatially decomposed visual evidence.

Conclusion

This study addressed the critical challenges in automated BIM Quality Assurance—specifically the "Semantic Gap" arising from linguistic ambiguities in LOD specifications and the "Geometric Complexity" caused by fragmentation and occlusion in 3D models. To bridge these gaps, we proposed a novel hybrid assessment framework that integrates a Bridge knowledge graph with a Semantic Geometric Preprocessing pipeline, validated through a MLLM (Gemini 2.5 Pro).

The experimental results quantitatively demonstrate the superiority of the proposed "Deep Standards" approach over conventional text-based methods. By translating abstract human requirements into system-specific visual inspection items, our framework achieved an overall accuracy of 83.47%, marking a significant improvement of 32.46% compared to the baseline. Notably, the recall rate increased dramatically from 21.61% to 82.26%, proving that the integration of the knowledge graph effectively mitigates the issue of missing feature detection inherent in traditional approaches. A four-condition ablation study further revealed that neither semantic decomposition (KG Only: 67.34%) nor geometric preprocessing (Cluster IMG Only: 48.19%) alone achieves comparable performance—confirming that the two components exhibit a complementary synergy, where the knowledge graph provides "what to inspect" and cluster images provide "where to inspect." The sensitivity analysis further confirmed that a feature Emergence threshold of 40% offers the optimal balance between precision and recall, ensuring robust performance even with complex geometric variations.

The qualitative contribution of this research lies in its transition from "Black-box" AI to an "Explainable" verification system. Unlike purely vision-based deep learning models that require massive labeled datasets, our system generates traceable rationales by cross-referencing logical criteria with semantic visual evidence. The implementation of "Semantic Units"—reassembled from fragmented meshes and visualized through material-based layer separation—enables the MLLM to accurately

identify internal components previously obscured by occlusion.

In conclusion, this framework successfully automates the translation of "Human Standards" into machine-interpretable "Deep Standards," providing a scalable and transparent solution for LOD verification. Future work will focus on expanding the dataset through systematic LOD augmentation—selectively adding or removing sub-components from existing high-fidelity families to generate boundary-condition models at each LOD transition—thereby enabling more rigorous evaluation of generalizability. Additionally, the "Visual Vocabulary" of the knowledge graph will be extended to cover a broader range of construction disciplines, and real-time inference speed will be optimized for large-scale infrastructure projects. This research paves the way for the next generation of intelligent BIM auditing tools, where domain knowledge and visual reasoning operate in synergy.

Acknowledgments

This work was supported by the National Research Foundation of Korea (NRF) grant funded by the Korea government (MSIT) (RS-2025-23324132) and the Korea Agency for Infrastructure Technology Advancement (KAIA) grant funded by the Ministry of Land, Infrastructure and Transport (RS-2024-00407028).

References

- Abualdenien, J. and Borrmann, A. (2022) 'Ensemble-learning approach for the classification of levels of geometry (LOG) of building elements', *Advanced Engineering Informatics*, 51, p. 101497.
- F.h, A., A.A., V.v, T. and J.h.m, T. (2025) 'BIM ontology for information management (BIM-OIM)', *Journal of Building Engineering*, 107, p. 112762.
- Nahri, I., Pinquicé, R., Véron, P., Bus, N. and Thorel, M. (2025) 'Extracting structured requirements from unstructured building technical specifications for building information modeling', *arXiv*.
- Utkucu, D., Ying, H., Wang, Z. and Sacks, R. (2024) 'Classification of architectural and MEP BIM objects for building performance evaluation', *Advanced Engineering Informatics*, 61, p. 102503.
- Wang, X., Yue, Q. and Liu, X. (2025) 'Crack image classification and information extraction in steel bridges using multimodal large language models', *Automation in Construction*, 171, p. 105995.
- Wu, J. and Zhang, J. (2019) 'New automated BIM object classification method to support BIM interoperability', *Journal of Computing in Civil Engineering*, 33(5), p. 04019033.
- Yang, M., Hei, X., Meng, H., Chen, K., Tong, X., Li, Y. and Zhao, Q. (2026) 'Lightweight framework for IFC element classification using multi-view and heterogeneous graph with decision-level fusion', *Automation in Construction*, 181, p. 106660.

- Yin, M., Tang, L., Webster, C., Xu, S., Li, X. and Ying, H. (2023) 'An ontology-aided, natural language-based approach for multi-constraint BIM model querying', *Journal of Building Engineering*, 76, p. 107066.
- Yu, Y.S., Kim, S.H., Lee, W.B. and Koo, B.S. (2023) 'Ensemble-based deep learning approach for performance improvement of BIM element classification', *KSCE Journal of Civil Engineering*, 27(5), pp. 1898–1915.
- Zeng, Z., Goo, J.M., Wang, X., Chi, B., Wang, M. and Boehm, J. (2024) 'Zero-shot building age classification from facade image using GPT-4', *arXiv*.