

Article

An Ontological Framework for Multidimensional and Multivariate Data Visualization with Applications to Financial and Accounting Data

Snezana Savoska ^{1,*}  and Suzana Loshkovska ^{2,*} ¹ University St. Kliment Ohridski Bitola, Faculty of Information and Communication Technologies, 7000 Bitola, North Macedonia² Ss. Cyril and Methodius University, Faculty of Computer Science and Engineering, 1000 Skopje, North Macedonia

* Correspondence: snezana.savoska@uklo.edu.mk (S.S.); suzana.loshkovska@finki.ukim.mk (S.L.)

Abstract

Selecting an appropriate visualization technique for multidimensional and multivariate financial and accounting (F&A) data remains a complex, user-dependent task. The TaxUI&BV4FADA taxonomy previously organized this problem along four dimensions—visualization techniques, user intentions and analytical goals, interaction possibilities, and user groups—but as a human-readable structure, it could not be queried, validated, or integrated into semantic decision-support pipelines. This paper presents an ontological framework that extends TaxUI&BV4FADA into a machine-readable OWL DL artifact authored in WebProtégé, with OWL used for semantic structuring and SPARQL used for score-based recommendation retrieval. The framework formalizes the four taxonomy dimensions and adds a decision-support layer and an evaluation layer. An explicit F&A semantic mapping is provided, and two contrasting worked scenarios—a financial analyst testing a gross-margin hypothesis and a CFO seeking a quarterly overview—show that the framework discriminates between F&A roles and analytical tasks. The evaluation demonstrates logical consistency, competency-question satisfaction, and internal consistency of the populated recommendation matrix against taxonomy-derived expectations, rather than independent empirical recommendation accuracy. This constitutes an internal, artifact-centered validation rather than an external empirical study with end users, and a protocol for future empirical validation with financial and accounting professionals is outlined. The framework provides a domain-oriented semantic and matrix-based decision-support foundation on which executable F&A visualization recommenders can be built.

Keywords: ontology; visualization; multidimensional and multivariate data; financial and accounting data; user intention; decision support; recommendation matrix; internal evaluation; SPARQL



Academic Editor: Olga Kurasova

Received: 17 June 2026

Revised: 16 August 2026

Accepted: 18 August 2026

Published: 20 August 2026

Copyright: © 2026 by the authors.

Licensee MDPI, Basel, Switzerland.

This article is an open access article distributed under the terms and conditions of the [Creative Commons Attribution \(CC BY\)](https://creativecommons.org/licenses/by/4.0/) license.

1. Introduction

Ontologies provide formal, machine-readable representations of domain concepts and their interrelationships [1,2]. In knowledge-based systems, they function as structured vocabularies that provide explicit definitions of domain entities, properties, and relations, thereby supporting knowledge sharing, reuse, interoperability, and reasoning over domain knowledge [3,4]. When combined with formal reasoning mechanisms, inference rules, and application-specific decision criteria, ontologies can further contribute to knowledge

inference and automated decision support. By reducing semantic ambiguity and enabling consistent data interpretation, ontologies support the systematic evaluation of alternative solutions, which makes them particularly suitable for complex, multi-criteria decision processes such as the selection of an appropriate visualization technique.

Visualization has become a fundamental approach for analysing, validating, and comparing data across business, scientific, and technological domains. The increasing volume and complexity of multidimensional and multivariate (MDMV) data have led to a wide spectrum of visualization techniques, including parallel coordinates [5], scatterplot matrices [6], glyph-based representations [7,8], dimensional stacking [9], pixel-oriented displays [10], dashboards [11], and hierarchical visualizations [12,13]. The selection of an appropriate technique depends not only on data characteristics but also on user-related factors, such as analytical intent, expertise, and interaction requirements [14–16]. This complexity is particularly pronounced in F&A analytics, where heterogeneous managerial roles and decision contexts interact with diverse visualization needs.

A substantial body of research has addressed the classification of visualization techniques through taxonomies based on data types, visual structures, tasks, or implementation aspects [17]. While these frameworks effectively organize the design space, they generally overlook user-centric and contextual factors and provide limited support for automated, user-driven selection. To address this gap, we previously introduced TaxUI&BV4FADA [18,19], a taxonomy tailored to MDMV visualization in F&A analysis. The framework structures selection criteria around user groups, analytical intentions, expected visualization effects, and interaction capabilities. Although effective as a conceptual guide, it remains a human-readable artifact and cannot be queried, validated, or integrated into semantic decision-support pipelines.

In this paper, we extend TaxUI&BV4FADA into an ontological framework that enables machine-readable representation, semantic querying, and score-based recommendation retrieval. The ontology is implemented in OWL DL and RDF [20] using the WebProtégé platform [21]. It formalizes the four taxonomy dimensions through four top-level classes—visualization techniques, classification criteria, interaction possibilities, and user groups—and adds a decision-support layer together with an evaluation layer. The current artifact contains a complete ABox population with respect to the $3 \times 4 \times 13$ technique-intention-user-group suitability matrix, resulting in 156 weight entries and the corresponding 156 ranked recommendations. The OWL ontology stores the decision structure and the recommendation assertions, while the numerical ranking is retrieved from the stored ABox scores through SPARQL queries rather than inferred by OWL DL reasoning. An explicit F&A semantic mapping documents the domain reading of the ontology elements and supports interpretation of the recommendation contexts. Two contrasting worked scenarios show that the framework discriminates between F&A roles and analytical tasks rather than producing a single generic answer. The ontology is positioned as a domain-oriented semantic and mathematical foundation on top of which an application-level recommender for live F&A systems can be built.

The paper makes four contributions: (i) it formalizes TaxUI&BV4FADA as a finite OWL DL ontology, with 198 classes organized around the four taxonomy-derived dimensions and extended by a decision-support and an evaluation layer; (ii) it defines a decision-support model over typed sets of techniques, intentions, and user groups, equipped with a suitability matrix $M : \mathcal{T} \times \mathcal{U} \times \mathcal{G} \rightarrow [0, 1]$ and a score-induced ranking; (iii) it materializes the full $3 \times 4 \times 13$ suitability matrix in the ABox together with the corresponding 156 ranked recommendation individuals spanning all twelve (intention, user-group) decision contexts; and (iv) it validates the artifact's internal consistency through structural metrics and

eight SPARQL competency queries, and illustrates the resulting recommendation behavior through two contrasting F&A scenarios.

It is important to delimit the scope of this contribution. The present work is the formalization of visualization-selection knowledge as a reusable, machine-readable semantic artifact, not the implementation of a complete, deployed recommendation system. The artifact provides the semantic foundation—classes, properties, a populated weight matrix, and SPARQL retrieval—on which such an operational recommender can subsequently be built, validated with end users, and integrated into financial analytics environments.

The remainder of the paper is organized as follows. Section 2 reviews related work on visualization taxonomies, visualization ontologies, and ontology-driven approaches in finance and accounting, and positions the present framework. Section 3 describes the methodology used to derive and validate the ontology. Section 4 presents the formal structure of the ontology, its decision-support and evaluation layers, the F&A semantic mapping, and the two worked recommendation scenarios. Section 5 reports the structural metrics, ABox retrieval functionality, and internal consistency and implementation-fidelity results. Section 6 discusses the scope of the contribution, the internal-versus-external validation strategy, and the practical applicability, scalability, and maintenance of the framework. Section 7 concludes, and Section 8 sets out directions for future work.

2. Related Work

The selection of appropriate visualization techniques for multidimensional and multivariate (MDMV) data depends on multiple interrelated factors, including data characteristics, analytical tasks, and user-related aspects such as intentions, expertise, and interaction requirements. Research addressing this problem has evolved along two main directions: (i) taxonomies and classification frameworks that organize visualization techniques, and (ii) ontological approaches that formalize visualization knowledge into machine-readable representations.

2.1. Visualization Technique Classification

The systematic classification of visualization techniques originates from foundational work on graphical representation and perception. Bertin's semiology of graphics [22] introduced visual variables and their perceptual properties, forming the basis for later classification frameworks. McCormick et al. [23] emphasized the need for structured organization of visualization methods, while Tufte [24] established principles for effective quantitative visualization.

Subsequent work focused on operational and design-oriented taxonomies. Shneiderman [14] proposed a task-by-data-type taxonomy linking visualization techniques to analytical tasks, while Card et al. [15] formalized the visualization process as a pipeline of data transformations, visual mappings, and interaction techniques. Mackinlay [25] introduced formal criteria for evaluating visual encodings, and later approaches [26,27] emphasized interaction and abstraction.

Although these taxonomies provide valuable insights into the visualization design space, they primarily focus on data properties, visual structures, and tasks. User-related aspects such as analytical intention, expertise, and interaction context are not systematically integrated. Consequently, these approaches remain descriptive and require additional formalization before they can support automated or adaptive visualization selection.

2.2. Visualization Ontologies

Research at the intersection of visualization and ontology splits into two distinct strands: visualization of ontologies, which develops graphical methods for exploring and

interpreting semantic structures [28], and ontologies for visualization, which use ontologies to model the visualization process itself. The present work belongs to the second strand; the first is mentioned only for disambiguation.

Within the ontologies-for-visualization line, foundational work by Duke et al. [29] introduced a shared conceptual structure that organizes visualization knowledge into data, tasks, processes, and representations, and Shu et al. [17] extended this direction with semantic descriptions for visualization services. More recent contributions narrow the scope to specific aspects of the visualization process: VISO [30] provides a structured representation of visualization components, mappings, and effectiveness knowledge; VisKo [31] formalizes workflow composition and service-based visualization pipelines; VIS4ML [32] targets visual analytics for machine-learning workflows; and AAVO [33] addresses audio-visual analytics.

A comprehensive survey of this body of work [34] shows that, despite a decade of progress, existing visualization ontologies remain fragmented, heterogeneous, and inconsistent in terminology and modeling strategies. Most encode only an isolated slice of the visualization process—pipelines, techniques, or services—without integrating data, visual encodings, tasks, interaction, and user context within a single representation.

A parallel thread emphasizes the role of analytical intent in visualization selection. Larrea et al. [16] demonstrate that visualization effectiveness depends strongly on the analytical goal: different tasks demand different visual representations even for identical datasets. Although some ontologies include tasks or user expertise, they rarely connect user roles, analytical intentions, and interaction requirements to an explicit recommendation mechanism, and the financial and accounting setting is essentially absent from this literature.

Two specific gaps, therefore, remain. First, coverage of the visualization pipeline is uneven, with current ontologies favoring particular slices of the visualization process, such as representation components, services, or workflows. Second, user-centered selection drivers—intentions, roles, and interaction requirements—are not linked to an explicit recommendation mechanism. Together, these gaps significantly restrict the applicability of existing ontologies in intelligent, user-adaptive visualization systems, and they motivate the framework proposed in this paper.

2.3. Ontology-Driven Visualization in Finance and Accounting

In the F&A domain, existing work is not only limited in number but also organized around specific systems, indicators, fraud-detection tasks, or reporting infrastructures. It also intersects with two practice traditions that predate the ontology work itself: the F&A dashboard/KPI literature [35,36] and the broader business-intelligence literature [37], both of which establish that financial reporting is multidimensional, role-conditioned, and dashboard-centric. The ontology-driven contributions themselves fall into four clusters; each provides a relevant point of comparison, and each leaves a gap addressed by the present framework.

The most developed cluster is the InKoM line (from the Polish *Inteligentny Kokpit Menedżerski*, “Intelligent Managerial Dashboard”), which uses financial ontologies to structure intelligent dashboards for small and medium-sized enterprises [38–41]. In InKoM, ontologies are not used merely for data annotation but actively shape the visual interface, enabling semantic navigation, contextual explanation of financial indicators, and guided analytical exploration.

However, InKoM operates at the level of a single dashboard configuration per enterprise: the link between visualization techniques, user intentions, and managerial roles is embedded implicitly inside each deployed dashboard rather than represented as a first-

class, queryable structure. As a result, the selection logic that determines which technique is appropriate for a given (intention, user-group) context is not formalized and cannot be inspected or reused independently of a particular enterprise deployment. The remaining three clusters reproduce the same structural limitation in different forms; in what follows, we therefore identify only their distinctive contributions and the distinctive gaps that the present framework addresses.

A second cluster within the same research line uses visualization for semantic explanation of financial indicators [42]. Its distinctive contribution is the use of semantic networks and topic maps to expose dependencies among indicators such as ROI, liquidity, and profitability, so that visualization serves not only as a presentation tool but also as a reasoning aid. The distinctive gap is that the locus of the model is the indicator-dependency graph rather than the technique-selection step: no mapping is provided from an analytical reasoning need to an appropriate visualization technique, so the choice of technique remains outside the formal model.

A third cluster introduces limited forms of user adaptation through user-knowledge profiles and process-oriented models [43]. Its distinctive contribution is the modeling of expertise and behavior through a novice/expert distinction and predefined analytical workflows. The distinctive gap is that user intentions and analytical goals are not represented as first-class ontology elements: the workflows adapt to whom the user is but not to what the user is currently trying to achieve.

Outside this main research line, the literature is sparse and case-specific. A fourth strand applies ontology-driven visualization to financial fraud detection through graph-based representations of transactions and relationships [44,45]; its distinctive contribution is the encoding of rich inter-account structure that supports a fixed visual idiom—network views—and its distinctive gap is task-specificity, since the encoded knowledge does not generalize beyond the fraud-detection slice of the F&A reporting cycle. A related strand focuses on XBRL and linked-data infrastructures, in which ontologies are used to structure financial reporting data and to support browsing and integration interfaces [46–49]; its distinctive contribution is the linked-data representation of the underlying reporting facts, and its distinctive gap is that these works sit on the data-infrastructure side of the problem—the ontology formalizes what is reported, but not how it should be visualized, so technique selection is left to the consuming interface.

Taken together, these four clusters show that ontology-driven visualization in finance and accounting is feasible, but they address different and relatively narrow parts of the problem. InKoM demonstrates the value of ontology-supported dashboards; semantic-indicator approaches show how financial concepts can be explained visually; fraud-detection systems illustrate task-specific graph-based analysis; and XBRL/linked-data approaches provide semantic infrastructure for financial reporting data. However, none of these lines exposes a reusable visualization-selection model in which user group, analytical intention, interaction requirement, candidate technique, suitability score, and rank are represented as first-class, queryable ontology elements. This is the specific gap addressed by the proposed framework.

2.4. Gap Analysis and Positioning

The reviewed literature reveals four persistent gaps. First, visualization taxonomies provide useful conceptual classifications but lack formal semantics and cannot be queried directly. Second, existing visualization ontologies provide machine-readable representations but usually focus on isolated aspects of the visualization process, such as services, workflows, visual mappings, or representation components. Third, user-centered selection factors—especially user roles, analytical intentions, and interaction requirements—are rarely represented as first-class elements of an explicit recommendation mechanism. Fourth,

ontology-driven visualization in finance and accounting remains mostly system-specific, focusing on dashboards, indicator explanation, fraud detection, or reporting-data infrastructure rather than on reusable technique recommendation. Table 1 summarizes the positioning of representative approaches.

To address these limitations, this paper formalizes the TaxUI&BV4FADA taxonomy as an OWL-based decision-support ontology. The proposed framework represents visualization techniques, user-centered classification criteria, interaction possibilities, and user groups within a unified semantic model, and adds a decision-support layer in which decision contexts, candidate techniques, suitability scores, and ranks are explicitly represented as queryable ontology elements.

Table 1. Positioning of representative visualization classification and ontology-based approaches. ✓ explicitly supported; △ partly/implicitly supported; – not supported.

Framework	Type	Main Focus	Techniques	Interaction	User Roles	User-Centered Selection	Recommendation Mechanism
General visualization taxonomies and ontologies							
Bertin [22]	Theory	Visual variables	△	–	–	–	–
Shneiderman [14]	Taxonomy	Tasks and data types	✓	✓	–	△	△
Card et al. [15]	Model	Visualization pipeline	✓	✓	–	–	–
Duke et al. [29]	Ontology	Visualization concepts	✓	△	–	–	–
Shu et al. [17]	Ontology	Visualization services	✓	△	–	–	△
VISO [30]	Ontology	Visualization representation	✓	△	–	–	△
VisKo [31]	Ontology	Visualization workflows	✓	△	–	–	△
VIS4ML [32]	Ontology	ML visualization workflows	✓	△	–	–	–
AAVO [33]	Ontology	Audiovisual analytics	✓	△	–	–	–
Larrea et al. [16]	Semantic approach	Task-dependent visualization	✓	–	–	△	△
Finance and accounting domain approaches							
Korczak & Dudycz [38–40]	Ontology system	Financial dashboards	✓	✓	△	△	△
Dudycz [42]	Ontology-based vis.	Semantic financial indicators	△	△	△	△	–
Fraud analysis [44,45]	Ontology application	Financial fraud detection	△	△	–	△	–
XBRL/Linked Data [46–49]	Semantic data systems	Financial reporting visualization	△	△	–	–	–
TaxUI&BV4FADA [18]	Taxonomy	User-centered selection	✓	✓	✓	✓	△

3. Materials and Methods

This section describes the methodology used to derive the proposed ontology from the TaxUI&BV4FADA taxonomy [18,19]. The aim is to specify the engineering steps, the tooling, the construction of the weight matrix, and the validation strategy; the resulting OWL artifact (its classes, properties, individuals, decision-support layer, and evaluation layer) is described in detail in Section 4.

3.1. Methodological Approach

The work follows a design-science approach in which an existing conceptual artifact, the TaxUI&BV4FADA taxonomy, is formalized, extended, and operationalized into a computational artifact [50,51]. The taxonomy serves as the requirements input; the OWL ontology and the decision-support layer described in Section 4 serve as the constructed artifact.

The ontology engineering steps follow the Noy and McGuinness methodology [3], in particular the iterative sequence of scope definition, term enumeration, class hierarchy construction, property definition, and instance population. The process was further informed by the W3C OWL and RDF recommendations [20,52] and by the WebProtégé collaborative authoring environment [21]. The eight phases adopted in the present work are summarized in Table 2.

Table 2. Methodological phases adopted in the present work, following the Noy and McGuinness ontology engineering methodology under a design-science framing.

Phase	Activity	Purpose
P1	Scope definition	Delimit the application domain and the boundaries of the ontology.
P2	Competency-question elicitation	Specify the queries the ontology must support and the structural properties it must satisfy.
P3	Knowledge-source analysis	Map the TaxUI&BV4FADA taxonomy elements onto candidate ontology elements.
P4	Conceptual modeling	Distinguish entities (classes), relationships (object properties), and attributes (data properties).
P5	Implementation	Encode the ontology in OWL DL using WebProtégé and serialize it in RDF/XML.
P6	Decision-support extension	Add the constructs required to represent decision contexts, recommendations, and weight evidence.
P7	Instance population	Populate the ABox with individuals corresponding to each concrete decision case, including the construction of the suitability matrix.
P8	Validation	Apply structural, competency-question-based, and internal-consistency checks.

The artifact-level results of phases P3–P7 (taxonomy-derived classes and subclasses, the four taxonomy superclasses, the decision-support classes, the evaluation classes, the populated weight matrix, and the ranked recommendation set) are reported in Section 4. Phase P8 is summarized methodologically below (Section 3.7); the corresponding quantitative results are reported in Sections 5.3 and 5.4.

3.2. Scope Definition

The scope of the ontology was defined to correspond directly to the TaxUI&BV4FADA taxonomy, ensuring a consistent mapping between the conceptual classification framework and its formal semantic representation. The ontology is positioned as a domain-specific, decision-support ontology focused on the selection of MDMV visualization techniques for F&A data. The taxonomy organizes the visualization-selection problem into four interrelated dimensions: visualization techniques, user intentions and analytical goals, interaction possibilities, and user groups; these dimensions delimit the boundaries of the ontology and guide the identification of core concepts. Low-level graphical rendering, perceptual modeling at the psychophysical level, and system-specific implementation details are explicitly out of scope unless they directly influence the selection of a visualization technique.

A secondary objective is extensibility and reuse. The ontology is designed as a modular, evolvable structure that can accommodate additional concepts, domains, and decision criteria without modifying the existing taxonomy backbone.

3.3. Competency Questions

Competency questions were used to define the functional requirements of the ontology and to guide its design [3]. They specify the types of queries the ontology must answer and

serve as a validation mechanism for completeness. The competency questions cover six facets that together exercise the decision-support and evaluation layers of the artifact:

- single-best-answer retrieval (which technique is recommended for a given user group and analytical intention?);
- top-*k* retrieval (what are the top three techniques for a given context?);
- interaction-aware retrieval (which interaction mode is considered in a given context?);
- threshold-based retrieval (which techniques achieve a recommendation score above a given threshold for a given context?);
- full-ranking retrieval (what is the complete ranked list, with scores and ranks, for a given context?);
- structural and completeness retrieval (how many weight-matrix entries are encoded in the full matrix, and does every decision context have exactly thirteen candidate recommendations?).

These six facets are operationalized through eight competency questions, encoded as named individuals CQ1–CQ8 in the OWL file. Two facets carry more than one question to widen coverage: single-best-answer retrieval is instantiated in two contrasting contexts (top managers seeking a data overview, CQ1; analytical staff performing hypothesis confirmation, CQ2), and the structural/completeness facet combines a global matrix-cardinality check (CQ5) with a per-context completeness check (CQ8). The remaining four facets are each instantiated by a single question—top-*k* retrieval as CQ3, interaction-aware retrieval as CQ4, threshold-based retrieval as CQ6, and full-ranking retrieval as CQ7.

Each competency question is mapped to a SPARQL query and stored inside the ontology together with the question's natural-language text, expected answer, and Boolean validation outcome. This design makes the validation layer self-describing: the artifact carries not only the domain knowledge but also a minimal executable validation protocol, so the same artifact can be re-validated after any ABox update without external bookkeeping or out-of-band query files. The eight encoded competency questions and their outcomes are reported in Section 5.3.

3.4. Conceptual Modeling

The taxonomy was interpreted as a structured conceptual model in which visualization techniques are related to user intentions, interaction capabilities, and user profiles. The conceptual model distinguishes between entities (modeled as OWL classes), relationships (modeled as object properties), and attributes (modeled as data properties); this abstract mapping between the taxonomy's modeling primitives and OWL primitives is operationalized in Step 2 of the methodological pipeline shown in Figure 1. The same distinction is preserved in the implemented OWL artifact and is used in Section 4 to introduce the formal vocabulary.

Following the Noy and McGuinness sequence, the four taxonomy-derived top-level dimensions of the visualization-selection problem were first identified as TBox seeds; subsequent phases refined those classes into a subclass hierarchy, attached object and data properties, populated the ABox with named individuals, and finally added the decision-support and evaluation extensions on top of the taxonomy-derived backbone. Figure 1 summarizes this construction as an eight-phase pipeline.

While Figure 1 captures the process of construction, Figure 2 captures the corresponding content-level mapping from the original TaxUI&BV4FADA taxonomy dimensions to the implemented OWL ontology. The four taxonomy-derived dimensions are formalized as the top-level ontology classes (MDMV_techniques; Tax4UIVB with its UserIntention branch; InteractionPoss; and UserGroups), while the decision-support and evaluation constructs are added as the two extensions required for recommendation and validation. The resulting

class hierarchy, the object and data properties, the populated ABox, and the decision-support and evaluation layers are presented in Sections 4.1–4.7 and Section 4.8, respectively.

Methodological Phases — Concrete Outputs	
Phase / activity	Key outputs (concrete results)
P1 Scope definition	Scope definition - F&A domain delimited · 4 user groups (UG_*) · 3 intentions (Int_*) 13 MDMV visualization techniques in scope
P2 Competency-question elicitation	Competency-question elicitation - CQ1–CQ8 formulated (structural + retrieval queries) - SPARQL validation queries VQ_CQ1–VQ_CQ8 specified
P3 Knowledge-source analysis	Knowledge-source analysis - TaxUI&BV4FADA taxonomy mapped → 198 candidate OWL classes 18 relations identified · 19 attributes identified - Entities, relations, and attributes distinguished by type
P4 Conceptual modelling	Conceptual modelling 198 OWL classes: C_tax 61 · C_dec 97 · C_eval 40 187 subclass axioms · depth ≤ 5 · no disjointness / equiv. axioms 18 object properties · 19 data properties · 4 annotation properties 3 SWRL rules · 12 SWRL atoms
P5 Implementation	Implementation - OWL DL ontology encoded in WebProtégé · serialized as RDF/XML 198 classes · 41 properties · 187 subclass axioms declared 3 SWRL rules (suitability aggregation, rank derivation, context matching)
P6 Decision-support extension	Decision-support extension - Classes added: DecisionContext · Recommendation · WeightMatrixEntry 12 decision contexts (CTX_*) · suitability matrix $T \times U \times G \rightarrow [0,1]$ 3 intentions × 4 user groups × 13 techniques = 156 WME_* entries 156 REC_* recommendation individuals with hasRank assertions
P7 Instance population	Instance population 390 named individuals declared in ABox - Tech_*(13) · Crit_*(3) · UG_*(4) · Int_*(5) · CTX_*(12) - WME_*(156) · REC_*(156) · EVAL_*(12) + CQ/VQ individuals - SPARQL-based retrieval (not OWL DL inference) - Note: named assets 404; owl:NamedIndividual 390
P8 Validation	Validation - All 8 CQs verified via SPARQL (VQ_CQ1–VQ_CQ8) - EvalRun_FullMatrix individual · 12 EVAL_* individuals - Structural, competency-question, and consistency checks passed

Figure 1. Methodological pipeline of the ontology construction: eight sequential phases from the input TaxUI&BV4FADA taxonomy to the populated decision-support and validation layers.

3.5. Implementation and Tooling

The ontology was implemented in WebProtégé [21] and serialized in RDF/XML in conformance with the W3C OWL and RDF recommendations [20,52]. The expressivity profile is OWL DL, which guarantees decidability of standard reasoning tasks and compatibility with mainstream reasoners [53–55]. Logical consistency was checked with the HermiT reasoner integrated in Protégé Desktop, with default settings; the reasoner reported no inconsistencies and no unsatisfiable classes for the populated artifact. Because the present release is primarily taxonomic and assertional, this consistency check verifies structural well-formedness rather than the satisfaction of non-trivial semantic constraints.

For the structural audit reported in Section 5.1, the ontology was encoded in OWL DL and serialized as RDF/XML, and structural counts were extracted from the named OWL classes, object properties, data properties, annotation properties, named individuals, subclass axioms, and ABox assertions. SPARQL queries against the populated ABox were used to operationalize the competency questions and to retrieve ranked recommendations; each query is stored inside the ontology as the queryText data-property value of the corresponding ValidationQuery individual (Section 5.3); the eight queries are reproduced in full in Appendix B and were executed directly against the RDF graph. All identifiers used in the OWL file are domain-neutral; their financial and accounting reading is made explicit in Section 4.6.

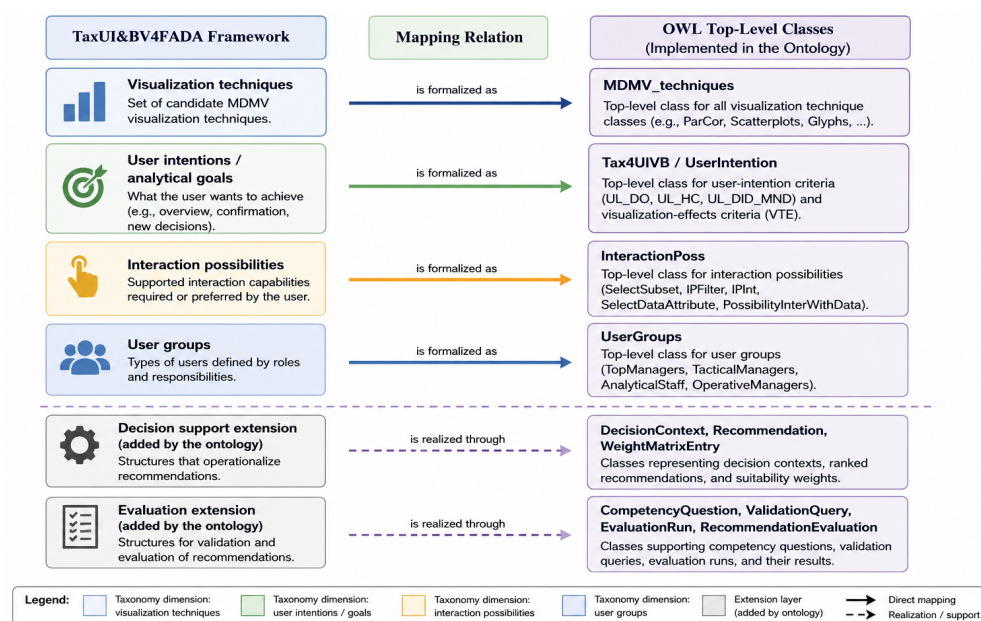


Figure 2. Mapping of the four TaxUI&BV4FADA taxonomy dimensions onto OWL top-level classes (top), together with the two extensions added in this work (bottom): the decision-support layer (C^{dec}) and the evaluation layer (C^{eval}). The individual class names appear in the body paragraph above and in Section 4.1.

3.6. Construction of the Suitability Matrix

The suitability matrix $M : \mathcal{T} \times \mathcal{U} \times \mathcal{G} \rightarrow [0, 1]$ that drives the recommendation layer was initialized heuristically from the TaxUI&BV4FADA taxonomy [19]. Each score $M(t, u, g) \in [0, 1]$ is an author-defined taxonomy consistency value that quantifies the suitability of technique t for an analyst with intention u and managerial role g . Concretely, the initial value $M(t, u, g)$ was set high, close to 0.99, when the taxonomy lists t as a recommended technique for role g under intention u , and progressively lower, down to 0.40, as the technique diverges from the taxonomy's expected ordering for that context. Scores were manually assigned by the authors according to this taxonomy-derived ordering; no fixed analytical rank-to-score formula, empirical usage data, or an independent expert-elicitation procedure, was used in the present release. The full $3 \times 4 \times 13 = 156$ matrix was then reified in the ABox through 156 WeightMatrixEntry individuals and stored on the matrixWeight data property.

The matrix is therefore not learned from empirical usage data and does not represent independent expert consensus elicited through a separate expert study. Its purpose at this stage is to provide a populated, queryable baseline that exposes the decision structure to SPARQL retrieval and that allows the framework to be tested for internal consistency against the taxonomy-derived expectations (Section 5.4). Replacing the taxonomy-

consistency initialization with expert-elicited weights is identified as the immediate next step beyond the present release (Section 4.11).

3.7. Validation Strategy

The ontology was validated at three complementary levels.

Structural validation. Logical consistency was checked with the HermiT reasoner, and the class hierarchy was audited along the standard structural dimensions (class, property, individual, axiom, and rule counts; hierarchy depth; branching factor). The audit and reasoner outcomes are reported in Sections 5.1 and 5.2.

Functional validation. Functional adequacy was assessed by expressing the competency questions as SPARQL queries and executing them against the populated ABox. Each competency question is reified in the OWL file as a *CompetencyQuestion* individual paired with a *ValidationQuery* individual. The questions, queries, and their outcomes are reported in Section 5.3.

Internal-consistency check of the recommendation layer. The rankings induced by the populated weight matrix were compared against the taxonomy-derived expected top alternatives. Because both the expected techniques and the matrix weights are seeded from the same taxonomy, this check measures implementation fidelity between the design specification and the populated ABox, not independent empirical accuracy. The measures used, their formal definitions, the storage pattern on *RecommendationEvaluation* individuals, and the per-context values are reported in Sections 4.8 and 5.4.

The three levels together exercise the ontology at the TBox level (structural metrics), at the ABox level (competency-question retrieval), and at the recommendation level (internal-consistency measures). Results are reported in Section 5.

4. Ontology for Taxonomy UI & BV4FADA

The ontology presented in this section formalizes the TaxUI&BV4FADA taxonomy of multidimensional and multivariate (MDMV) data visualization and extends it with a decision-support layer for recommending visualization techniques. The implementation is provided as an OWL DL model authored in WebProtégé. The artifact is organized around three logical groupings of classes—a taxonomy-derived backbone, a decision-support extension, and an evaluation layer—whose formal definition, partition, and cardinalities are given in Section 4.1. The current ontology is a framework: it provides the conceptual structure, the fully populated weight matrix, and a baseline evaluation layer, but it is not yet a fully integrated end-to-end recommender system.

Figure 3 gives an intuitive overview of this organization before the formal definition is introduced.

4.1. Formal Structure and Superclass Definition

Formally, the ontology is defined as the 10-tuple

$$O = \langle C, I, P_o, P_d, A_{\text{sub}}, A_{\text{type}}, A_{\text{obj}}, A_{\text{data}}, A_{\text{ann}}, R_{\text{swrl}} \rangle, \quad (1)$$

where C is a finite set of OWL classes, I a finite set of named individuals, and P_o, P_d finite sets of object and data properties. The TBox is represented by the class set C , the property vocabularies P_o and P_d , and the subsumption-axiom set $A_{\text{sub}} \subseteq C \times C$ (encoded as `rdfs:subClassOf`). The ABox is represented by the class-assertion set $A_{\text{type}} \subseteq I \times C$ (encoded as `rdf:type`) together with the property-assertion sets

$$A_{\text{obj}} \subseteq I \times P_o \times I, \quad A_{\text{data}} \subseteq I \times P_d \times \mathcal{L}, \quad (2)$$

with $\mathcal{L}$ the literal space (typed RDF literals in $\mathbb{R}$, $\mathbb{B}$, or Σ^*). The set A_{ann} collects annotation assertions and R_{SWRL} the SWRL rules.

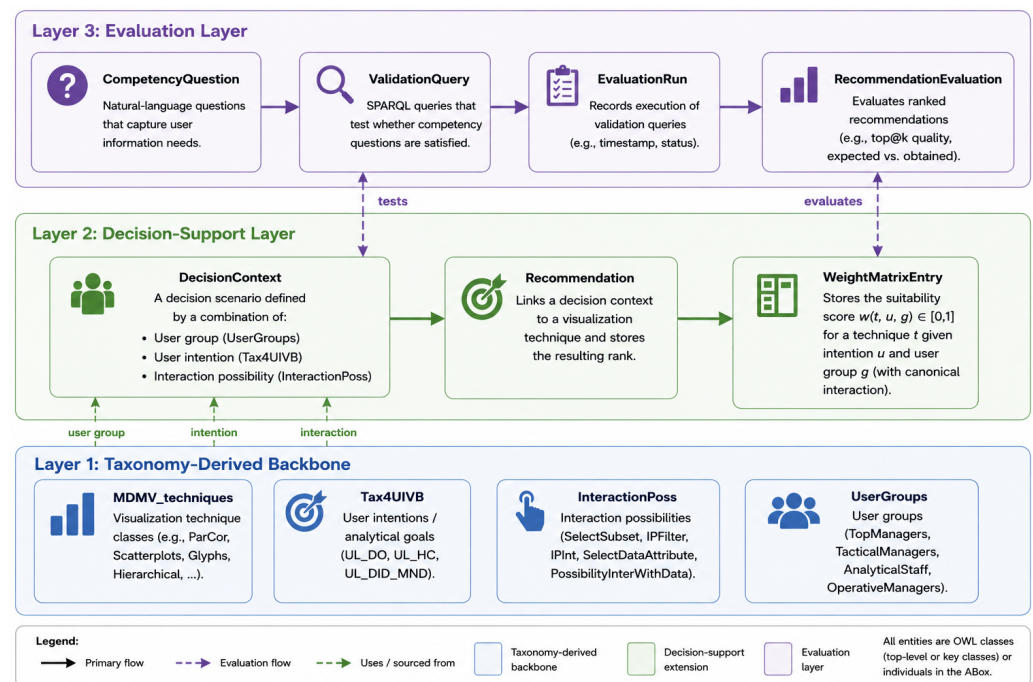


Figure 3. Layered architecture of the proposed ontology. (**Bottom**): taxonomy-derived backbone C^{tax} (MDMV_techniques, Tax4UIVB, InteractionPoss, UserGroups). (**Middle**): decision-support layer C^{dec} (see Section 4.7). (**Top**): evaluation layer C^{eval} (see Section 4.8).

Figure 4 provides an intuitive view of the formal distinction introduced above. The TBox contains the schema-level vocabulary of the ontology, including classes, object properties, data properties, subclass axioms, and declared domain and range axioms. The ABox contains named individuals and the concrete assertions used in the populated decision-support and evaluation layers. In the present ontology, these individuals include technique individuals, criterion and intention individuals, user-group individuals, interaction individuals, decision-context individuals, weight-matrix entries, recommendation individuals, and evaluation individuals. The three assertion types shown in the middle of the figure correspond to the formal components A_{type} , A_{obj} , and A_{data} : class assertions instantiate individuals as classes, object-property assertions connect individuals to other individuals, and data-property assertions connect individuals to literal values such as scores, ranks, query texts, and validation flags.

The class set is organized at the meta level into a taxonomy-derived layer C^{tax} , a decision-support layer C^{dec} , an evaluation layer C^{eval} , and a residual subclass material C^{sub} :

$$C = C^{\text{tax}} \cup C^{\text{dec}} \cup C^{\text{eval}} \cup C^{\text{sub}}, \quad (3)$$

with

- **Top-level taxonomy classes:**
 $C^{\text{tax}} = \{ \text{MDMV_techniques}, \text{Tax4UIVB}, \text{InteractionPoss}, \text{UserGroups} \},$
- **Decision-support classes:**
 $C^{\text{dec}} = \{ \text{DecisionContext}, \text{Recommendation}, \text{WeightMatrixEntry} \},$
- **Evaluation classes:**
 $C^{\text{eval}} = \{ \text{CompetencyQuestion}, \text{ValidationQuery}, \text{EvaluationRun}, \text{RecommendationEvaluation} \}.$

The decomposition in (3) is a meta-level organizational partition of the named classes and is not currently enforced by OWL disjointness axioms: the artifact contains zero `owl:disjointWith` and `AllDisjointClasses` axioms, although the four named-class sets are pairwise disjoint by inspection.

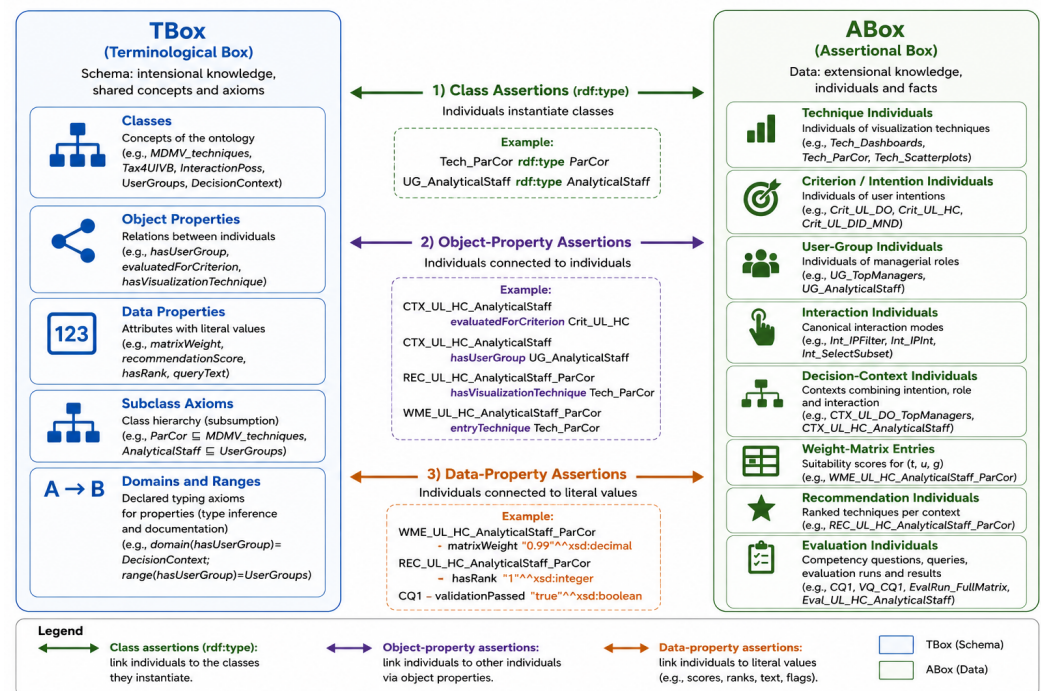


Figure 4. Schematic of the TBox–ABox split of the proposed ontology. The TBox (left) carries the schema-level vocabulary; the ABox (right) carries the populated named individuals; the three arrow styles in between denote the assertion sets A_{type} (class assertions, `rdf:type`), A_{obj} (object-property assertions), and A_{data} (data-property assertions on literal values).

The corresponding cardinalities are $|C^{tax}| = 4$, $|C^{dec}| = 3$, $|C^{eval}| = 4$, $|C| = 198$, so $|C^{sub}| = 187$. The subsumption set A_{sub} contains $|A_{sub}| = 187$ axioms, and the induced relation (C, A_{sub}) is a finite, acyclic class hierarchy rooted in `owl:Thing`, with eleven top-level named branches.

The ontology has eleven top-level named classes, i.e., classes whose only `rdfs:subClassOf` parent is `owl:Thing`— $|C^{tax}| + |C^{dec}| + |C^{eval}| = 4 + 3 + 4 = 11$. Within this top-level set, the four classes of C^{tax} form the taxonomy backbone of the ontology and serve as the roots of the refinement trees inside C^{sub} , whose internal structure is developed in Sections 4.2–4.5.

The property set splits into object and data properties, $|P_o| = 18$ and $|P_d| = 19$, together with four annotation properties. Each object property $p \in P_o$ carries stated domain and range classes $\text{dom}(p), \text{rng}(p) \in C$; the corresponding ABox triples in A_{obj} relate individuals $x \in \text{dom}(p)^I$ to individuals $y \in \text{rng}(p)^I$. Each data property $d \in P_d$ carries a domain class $\text{dom}(d) \in C$ and a datatype range $\text{rng}(d) \subseteq \mathcal{L}$, with $\text{rng}(d) \in \{\mathbb{R}, \mathbb{B}, \Sigma^*\}$ in the present artifact; data assertions take the form $(x, d, \ell) \in A_{data}$ with $x \in \text{dom}(d)^I$ and $\ell \in \text{rng}(d)$. For each $c \in C$, the local property signature is

$$P(c) = \{p \in P_o \cup P_d \mid \text{dom}(p) = c\} \subseteq P_o \cup P_d. \quad (4)$$

The stated domain and range axioms are used here as part of the ontology’s semantic documentation and type-inference structure. Under OWL’s open-world semantics, however, they do not function as closed-world validation constraints: an assertion outside an expected domain or range does not by itself produce a validation error, but may in-

stead lead to additional type inferences. Closed-world constraint checking is therefore outside the scope of the present OWL-only release and is identified as future work through SHACL validation.

The named-individual set decomposes by the principal naming scheme of the ABox,

$$I = I^{\text{Tech}} \cup I^{\text{Crit}} \cup I^{\text{UG}} \cup I^{\text{Int}} \cup I^{\text{CTX}} \cup I^{\text{WME}} \cup I^{\text{REC}} \cup I^{\text{EVAL}} \cup I^{\text{CQ\&VQ}} \cup \{\text{EvalRun_FullMatrix}\} \cup I^{\text{legacy}} \cup I^{\text{pun}}, \quad (5)$$

with cardinalities 13, 3, 4, 5, 12, 156, 156, 12, 16, 1 for the ten full-matrix subsets, $|I^{\text{legacy}}| = 4$ for the legacy individuals, and $|I^{\text{pun}}| = 8$ for the class-individual punning pairs, summing to $|I| = 390$.

The legacy subset I^{legacy} contains four individuals: two preliminary decision-context examples, DC_DO_TopManagers_Dashboards and DC_HC_AnalyticalStaff_ParCor, together with their two recommendation counterparts, REC_DO_TopManagers_Dashboards and REC_HC_AnalyticalStaff_ParCor. These individuals pre-date the full-matrix population. They are retained in the OWL file for historical reference but are excluded from the 12-context/156-recommendation full-matrix evaluation reported in Section 5.4; throughout this paper, I^{CTX} denotes the twelve full-matrix CTX_* contexts and I^{REC} the 156 full-matrix REC_UL_* recommendations. The SWRL component R_{SWRL} contains three rules with twelve atoms in total and is discussed in Section 4.11.

The artifact is connected to its ABox through a family of meta-level representative maps $\rho_T, \rho_U, \rho_G, \rho_J$, one per taxonomy-derived dimension. Each $\rho_\bullet$ pairs a TBox class with its canonical ABox individual; together they assemble into the partial reification map $\iota: C \rightarrow I$, which assigns to each class its representative when one exists, for example $\iota(\text{UL_D0}) = \text{Crit_UL_D0}$, $\iota(\text{UL_HC}) = \text{Crit_UL_HC}$, $\iota(\text{UL_DID_MND}) = \text{Crit_UL_DID_MND}$. These maps are external naming conventions; they do not assert OWL class-individual identity and are unrelated to OWL punning (the 8 punned pairs in I^{pun} arise from independent artifact-level constructs).

Figure 5 visualizes 11 immediate children of `owl:Thing` that belong to $C^{\text{tax}} \cup C^{\text{dec}} \cup C^{\text{eval}}$.

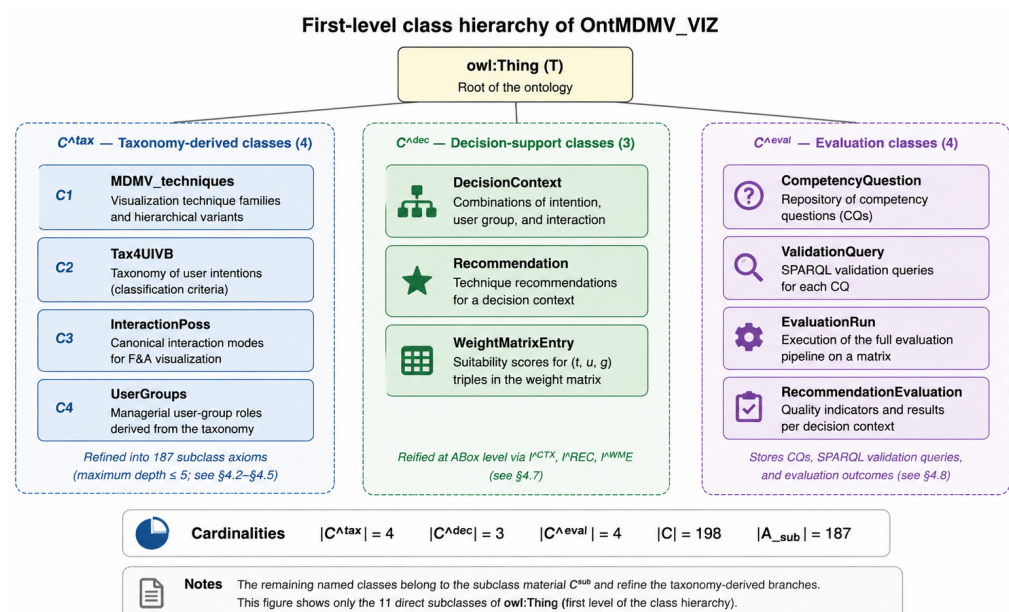


Figure 5. First-level class hierarchy of the ontology. The figure shows the eleven named classes whose direct superclass is `owl:Thing`, grouped into the taxonomy-derived layer C^{tax} , the decision-support layer C^{dec} , and the evaluation layer C^{eval} .

4.2. Superclass C_1 : Visualization Techniques (MDMV_Techniques)

The superclass $C_1 = \text{MDMV_techniques}$ forms the technique-oriented branch of O . We separate the TBox and ABox views explicitly. Let $T_C \subseteq C$ denote the set of technique classes obtained from the refinement tree $T(C_1) = \{c \in C \mid c \sqsubseteq^* C_1\}$, and let $T_I \subseteq I$ denote the corresponding set of technique individuals populated in the ABox; by construction $T_I = I^{\text{Tech}}$. The direct children of C_1 form the antichain

$$\begin{aligned} \text{Children}(C_1) = \{ & \text{Dashboards, Scatterplots, ParCor, Pixel_oriented_display,} \\ & \text{Glyphs, DimensionalStacking, Nodes_and_Links,} \\ & \text{Polar_coordinates} \}. \end{aligned} \quad (6)$$

The tree is enriched with the hierarchical-variant antichain

$$\begin{aligned} \text{Hier}(C_1) = \{ & \text{HierarParCor, HierarScatterplots, HierarPixOriented_display,} \\ & \text{HierarGlyphs, HierarDimenStacking} \}, \end{aligned} \quad (7)$$

and the dashboard branch is further refined by

$$\begin{aligned} \text{Sub}(\text{Dashboards}) = \{ & \text{histograms, scatter_plots, line_plots,} \\ & \text{bubble_plots, time_series} \}. \end{aligned} \quad (8)$$

The thirteen classes $T_C = \text{Children}(C_1) \cup \text{Hier}(C_1) \cup \text{Sub}(\text{Dashboards})$ are the technique categories reified at the ABox level by the meta-level representative map

$$\rho_T : T_C \xrightarrow{\sim} T_I, \quad |T_C| = |T_I| = 13, \quad (9)$$

with $\rho_T(\text{Dashboards}) = \text{Tech_Dashboards}$, $\rho_T(\text{ParCor}) = \text{Tech_ParCor}$, and so on. ρ_T is the restriction of ι to T_C ; it is an external naming convention and does not assert OWL class-individual identity.

The candidate set T_C intentionally combines visualization families, dashboard sub-techniques, and hierarchical variants, because it follows the granularity of the original TaxUI&BV4FADA taxonomy rather than a single uniform abstraction level. In particular, dashboards are treated as composite F&A visualization environments that may contain concrete chart types such as histograms, scatter plots, line plots, bubble plots, and time-series plots. Similarly, hierarchical variants, such as HierarParCor and HierarScatterplots are retained as separate candidates when the analytical task involves nested, grouped, or multi-level F&A structures, such as account hierarchies, product categories, organizational units, cost centers, or regional groupings. Thus, the thirteen elements of T_C should be interpreted as the practical recommendation candidates materialized in the ABox, not as a pure taxonomic level of equally abstract visualization families.

In contrast to a purely descriptive taxonomy, the ontology organizes these techniques into a formal class hierarchy in which each technique is a concept with explicitly defined properties and relationships. The decision-support layer reuses the technique individuals of T_I . For example, let r denote the recommendation individual for the high-cognitive-load analytical-staff context and the parallel-coordinates technique, and let w denote the corresponding weight-matrix entry. The relevant ABox assertions are

$$\begin{aligned} (r, \text{hasVisualizationTechnique}, \text{Tech_ParCor}) &\in A_{\text{obj}}, \\ (w, \text{entryTechnique}, \text{Tech_ParCor}) &\in A_{\text{obj}}, \\ (w, \text{matrixWeight}, m) &\in A_{\text{data}}, \quad m \in [0, 1]. \end{aligned} \quad (10)$$

The first assertion links the recommendation to the selected visualization technique, whereas the second and third assertions link the corresponding matrix entry to the same technique and to its suitability weight.

Figure 6 visualizes the operational candidate set introduced above and shows how the thirteen technique candidates used in the suitability matrix are obtained from the technique branch of the ontology. The figure makes explicit that the set T_C is not a single uniform taxonomic level, but a practical recommendation set inherited from $TaxUI\&BV4FADA$. It combines broad technique families, dashboard sub-techniques, and hierarchical variants, all of which are reified as corresponding $Tech_*$ individuals in the ABox and used in the $3 \times 4 \times 13$ suitability matrix.

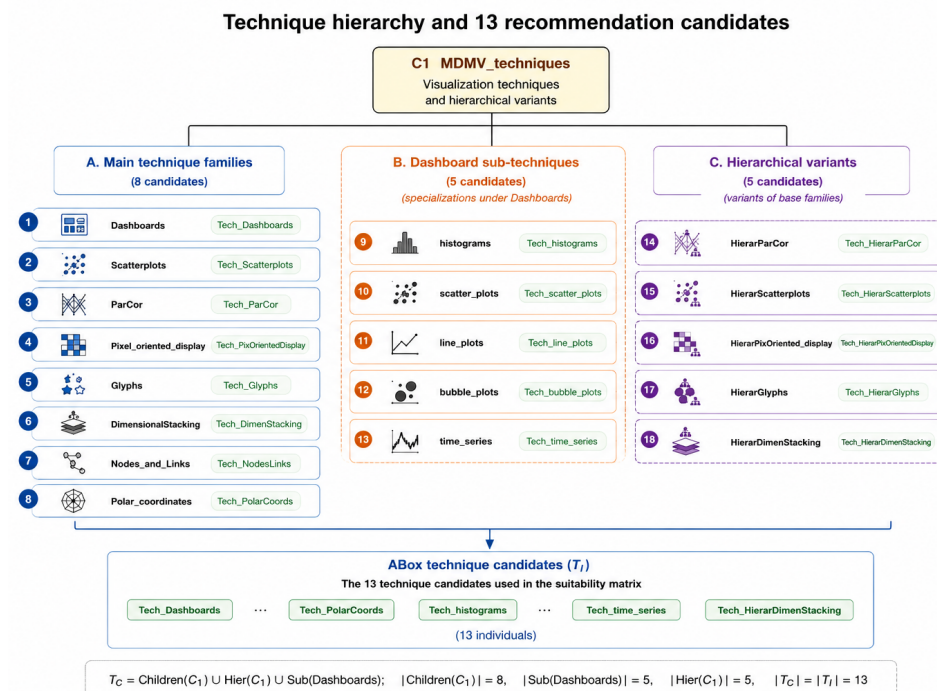


Figure 6. Technique branch of $C_1 = \text{MDMV_techniques}$ and the thirteen ABox candidates (T_C) used in the $3 \times 4 \times 13$ suitability matrix. Construction and rationale of T_C are described in the body paragraph above.

4.3. Superclass C_2 : Classification Criteria ($Tax4UIVB$)

The superclass $C_2 = \text{Tax4UIVB}$ models the classification criteria that relate techniques to user intentions and expected visualization effects. It decomposes as

$$C_2 = C_2^{UI} \cup C_2^{VTE}, \quad (11)$$

with $C_2^{UI} = \text{UserIntention}$ and $C_2^{VTE} = \text{VTE}$ at the class level.

User intention. The first component partitions into three subclasses,

$$U_C = \text{Children}(\text{UserIntention}) = \{UL_DO, UL_HC, UL_DID_MND\}, \quad (12)$$

with semantic readings

$UL_DO \mapsto$ Data Overview (exploration/summarization),

$UL_HC \mapsto$ Hypothesis Confirmation,

$UL_DID_MND \mapsto$ Delve Into Data and Make New Decisions.

The meta-level representative map for the user-intention dimension is

$$\rho_U : U_C \xrightarrow{\sim} U_I, \quad |U_C| = |U_I| = 3, \quad (13)$$

with image $U_I = I^{\text{Crit}} = \{\text{Crit_UL_DO}, \text{Crit_UL_HC}, \text{Crit_UL_DID_MND}\}$. Criterion individuals of U_I enter the decision-support layer through three object properties—hasClassificationCriterion, entryCriterion, and evaluatedForCriterion—whose ABox triples take the form $(x, p, y) \in A_{\text{obj}}$ with $y \in U_I$.

Visualization effects. The second component partitions into three first-level subclasses,

$$\text{Children}(\text{VTE}) = \{\text{Visibility}, \text{Interpretability}, \text{Insight_in_data}\}, \quad (14)$$

each further refined:

$$\begin{aligned} \text{Children}(\text{Visibility}) &= \{\text{VisAnAg}, \text{VisObjSel}, \text{VisRelation}, \text{VisScr}\}, \\ \text{Children}(\text{Interpretability}) &= \{\text{IntAgg}, \text{IntLevData}, \text{IntRelation}\}, \\ \text{Children}(\text{Insight_in_data}) &= \{\text{InsDataAss}, \text{InsDataClass}, \text{InsDataClu}, \text{InsDataCor}, \\ &\quad \text{InsDataPoss}, \text{InsDataPr}\}. \end{aligned}$$

These subclasses act as evaluative descriptors that support the qualitative interpretation of the taxonomy-derived weights stored in the recommendation matrix.

In the present release, however, the VTE branch is not used as an independent index of the suitability matrix. The matrix is defined over techniques, user intentions, and user groups, $M : \mathcal{T} \times \mathcal{U} \times \mathcal{G} \rightarrow [0, 1]$, while the VTE classes provide an explanatory vocabulary for interpreting why particular techniques are suitable in a given context. Thus, *Visibility*, *Interpretability*, and *Insight_in_data* document the taxonomy-derived rationale behind the score assignments, but they do not yet define a fourth matrix dimension.

Figure 7 visualizes the full hierarchy of C_2 . The left panel shows the user-intention branch $C_2^{UI} = \text{UserIntention}$, its three TBox subclasses, and their ABox images under ρ_U . The right panel shows the visualization effects branch $C_2^{VTE} = \text{VTE}$ with its three first-level subclasses and all thirteen leaf refinements. The VTE branch is included as explanatory vocabulary for interpreting taxonomy-derived suitability scores, whereas the matrix itself is indexed only by technique, intention, and user group.

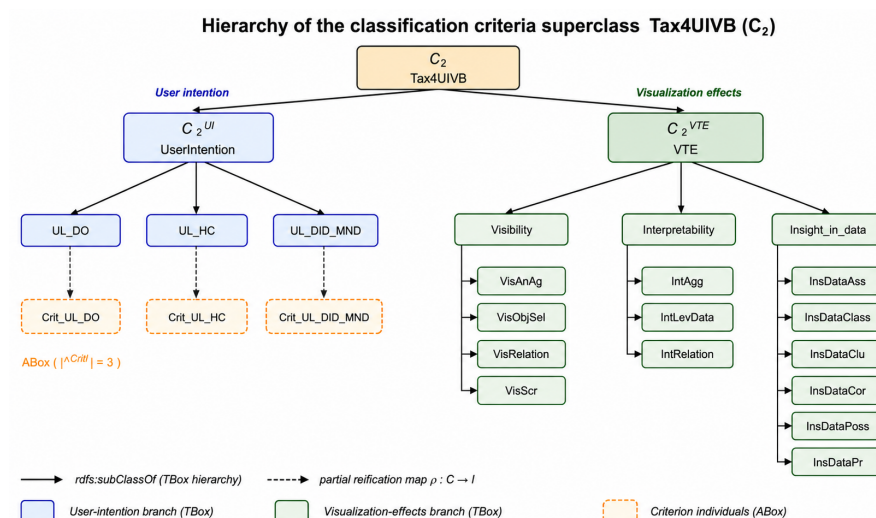


Figure 7. Hierarchy $C_2 = \text{Tax4UIVB}$. (Left): user-intention branch C_2^{UI} ; (Right): visualization-effects branch C_2^{VTE} . Solid arrows: rdfs:subClassOf ; dashed arrows: meta-level representative map ρ_U .

4.4. Superclass C_3 : Interaction Possibilities (*InteractionPoss*)

The superclass $C_3 = \text{InteractionPoss}$ captures interaction as an essential dimension. Its first-level partition is

$$J_C = \text{Children}(C_3) = \{ \text{SelectSubset}, \text{IPFilter}, \text{IPInt}, \text{SelectDataAttributes}, \text{PossibilityInterWithData} \},$$

with cardinality $|J_C| = 5$. Two of these classes are further refined:

$$\begin{aligned} \text{Children}(\text{IPFilter}) &= \{ \text{SFIL1}, \text{SFILn}, \text{SOBJn}, \text{STFIL} \}, \\ \text{Children}(\text{IPInt}) &= \{ \text{SDIS}, \text{SIF}, \text{SLB}, \text{SNO}, \text{SZOOM} \}. \end{aligned}$$

The interaction branch contains several compact identifiers inherited from the original taxonomy. To make their interpretation explicit, Table 3 provides a readable glossary for the refined interaction subclasses under IPFilter and IPInt. These subclasses are part of the TBox interaction vocabulary. In the current release, however, the populated decision contexts refer only to the five first-level interaction individuals listed below; the refined subclasses provide semantic detail for future extensions and for interpreting the interaction dimension.

Table 3. Readable interpretation of compact interaction identifiers in the *InteractionPoss* hierarchy. The identifiers are retained from the original taxonomy, while the readable labels clarify their semantic role in the ontology.

Parent Class	OWL Identifier	Readable Interpretation
IPFilter	SFIL1	Single filter operation
IPFilter	SFILn	Multiple filter operations
IPFilter	SOBJn	Selection of multiple objects
IPFilter	STFIL	Stepwise or temporal filtering
IPInt	SDIS	Display or visual-detail interaction
IPInt	SIF	Interactive filtering or focus operation
IPInt	SLB	Label-based interaction
IPInt	SNO	Node or object selection interaction
IPInt	SZOOM	Zooming interaction

The five first-level interaction classes are reified in the ABox through the meta-level representative map

$$\rho_J : J_C \xrightarrow{\sim} J_I, \quad |J_C| = |J_I| = 5, \quad (15)$$

with image

$$J_I = I^{\text{Int}} = \{ \text{Int_SelectSubset}, \text{Int_IPFilter}, \text{Int_IPInt}, \text{Int_SelectDataAttributes}, \text{Int_PossibilityInterWithData} \}.$$

Interaction individuals of J_I participate in the decision layer through two object properties—*considersInteraction* and *entryInteraction*—whose ABox triples $(x, p, y) \in A_{\text{obj}}$ have $y \in J_I$. Visualization techniques expose their supported interaction forms through the object property *hasInteractionPossibility*, whose assertions take the form $(t, \text{hasInteractionPossibility}, j) \in A_{\text{obj}}$ with t an instance of some class in T_C and j an instance of some class in J_C .

4.5. Superclass C_4 : User Groups (*UserGroups*)

The superclass $C_4 = \text{UserGroups}$ captures the heterogeneity of users targeted by the framework. Its first-level partition is

$$\begin{aligned} G_C &= \text{Children}(C_4) \\ &= \{ \text{TopManagers}, \text{TacticalManagers}, \\ &\quad \text{AnalyticalStaff}, \text{OperativeManagers} \}. \end{aligned} \quad (16)$$

with $|G_C| = 4$. Each user-group class is further refined into role-specific subclasses:

$$\begin{aligned} \text{Children}(\text{TopManagers}) &= \{ \text{PlanControl}, \text{RegMan}, \text{StratMan} \}, \\ \text{Children}(\text{TacticalManagers}) &= \{ \text{ManagByExep}, \text{SectorPC}, \text{TacOperLevel} \}, \\ \text{Children}(\text{AnalyticalStaff}) &= \{ \text{ExceptionAn}, \text{PerceptionAn}, \text{SpecAnTask} \}, \\ \text{Children}(\text{OperativeManagers}) &= \{ \text{OperManagers}, \text{ProbDetSolv}, \text{StProcControl} \}. \end{aligned}$$

These role-specific subclasses refine the semantic interpretation of the four primary user groups, but they are not used as separate indices of the suitability matrix in the present release. The recommendation matrix is populated at the level of the four primary user-group classes *TopManagers*, *TacticalManagers*, *AnalyticalStaff*, and *OperativeManagers*. The role-specific subclasses are therefore included as taxonomy-derived TBox refinements for documentation, explanation, and future extension, rather than as separate ABox decision categories in the current $3 \times 4 \times 13$ matrix.

Accordingly, the recommendation layer is populated at the primary level: the meta-level representative map is

$$\rho_G : G_C \xrightarrow{\sim} G_I, \quad |G_C| = |G_I| = 4, \quad (17)$$

with image

$$\begin{aligned} G_I = I^{\text{UG}} &= \{ \text{UG_TopManagers}, \text{UG_TacticalManagers}, \\ &\quad \text{UG_AnalyticalStaff}, \text{UG_OperativeManagers} \}. \end{aligned}$$

User-group individuals of G_I participate in the decision layer through three object properties—*hasUserGroup*, *entryUserGroup*, and *evaluatedForUserGroup*—whose ABox triples $(x, p, y) \in A_{\text{obj}}$ have $y \in G_I$.

The four representative maps $\rho_T, \rho_U, \rho_G, \rho_I$ defined in Sections 4.2–4.5 are meta-level naming conventions: they pair each TBox class with its canonical ABox individual but do not assert OWL class–individual identity. The 8 punned class–individual pairs counted in I^{pun} (Equation (5)) arise from independent artifact-level constructs and are unrelated to these representative maps.

4.6. Financial and Accounting Semantic Mapping

Let $\mathcal{D}_{\text{F\&A}}$ denote the codomain of financial and accounting concepts used as the target of the semantic mapping. We formalize $\mathcal{D}_{\text{F\&A}}$ as the disjoint union of five pairwise non-overlapping concept-sets,

$$\mathcal{D}_{\text{F\&A}} = D_{\text{role}} \dot{\cup} D_{\text{task}} \dot{\cup} D_{\text{interaction}} \dot{\cup} D_{\text{technique}} \dot{\cup} D_{\text{effect}}, \quad (18)$$

where D_{role} collects F&A managerial-role concepts, D_{task} the analytical-task concepts in the F&A decision cycle, $D_{\text{interaction}}$ the analyst-action concepts on F&A datasets, $D_{\text{technique}}$ the

F&A visualization-technique readings, and D_{effect} the perceptual-goal concepts on F&A reports. Each element of $\mathcal{D}_{\text{F\&A}}$ is a concept label rather than a formal OWL class.

The domain-neutral OWL identifiers used in Sections 4.2–4.5 are interpreted in this codomain through a semantic-mapping function

$$\sigma : C \cup I \longrightarrow \mathcal{P}(\mathcal{D}_{\text{F\&A}}), \quad (19)$$

that assigns each ontology element to the set of F&A concepts it denotes. The conceptual content of σ originates from the TaxUI&BV4FADA taxonomy [18], and σ is supplied extrinsically through Tables 4 and A1: it is not implemented as OWL axioms, `owl:imports`, or annotation assertions on the ontology elements. Consequently, σ supports explanation and domain interpretation but does not contribute to OWL entailment or to automated F&A-domain reasoning in the present artifact. Thus, the financial and accounting semantics of the present artifact are supplied as a documented interpretation layer over domain-neutral OWL identifiers, rather than through a separately imported F&A ontology. The ontology therefore formalizes the visualization-selection structure for F&A decision contexts, while the domain reading of roles, tasks, interactions, techniques, and visualization effects is made explicit through the semantic mapping σ . This design keeps the artifact lightweight and focused on technique recommendation, but it also means that deeper F&A-domain reasoning, such as reasoning over accounting standards, financial statement structures, or XBRL reporting facts, is outside the scope of the current release.

Two complementary tables document σ . Table 4; gives the compact dimension-level reading and is presented in this section. Appendix A contains the full element-level mapping in Table A1, with one row for each taxonomy-derived OWL class and ABox individual that participates in the recommendation layer. Both tables are interpretive: they document the F&A meaning of domain-neutral ontology identifiers but do not introduce additional OWL axioms or entailments.

For a decision context $\kappa \in I^{\text{CTX}}$, the F&A reading of κ is defined as the union of the F&A readings of its participating individuals, using the actual OWL properties that link a CTX_* individual to its components in the artifact:

$$\sigma(\kappa) = \sigma(ug) \cup \sigma(cr) \cup \sigma(int), \quad (20)$$

where ug, cr, int are the unique individuals satisfying

$$(\kappa, \text{evaluatedForUserGroup}, ug) \in A_{\text{obj}},$$

$$(\kappa, \text{evaluatedForCriterion}, cr) \in A_{\text{obj}},$$

$$(\kappa, \text{considersInteraction}, int) \in A_{\text{obj}}.$$

For instance, $\sigma(\text{CTX_UL_D0_TopManagers})$ encodes “CFO requires a quarterly data overview” and yields the top-ranked technique $\sigma(\text{Tech_Dashboards}) = \text{“F\&A KPI dashboard.”}$ Likewise, $\sigma(\text{CTX_UL_HC_AnalyticalStaff})$ encodes “financial analyst tests a cross-account hypothesis,” with top-ranked $\sigma(\text{Tech_ParCor}) = \text{“parallel-coordinates of financial ratios.”}$ The Maria scenario in Section 4.9 instantiates this second case. Because σ is external to the OWL axiomatization, the F&A readings above are documentation/interpretation artifacts; they support human explanation and downstream domain interpretation but are not consumed by any OWL DL reasoner in the present release.

Table 4. Dimension-level F&A reading of the semantic mapping σ . The OWL identifiers are domain-neutral; their F&A meaning is documented by σ rather than asserted as OWL axioms.

Ontology Dimension	F&A Interpretation	Typical Examples
Visualization techniques MDMV_techniques	Visual representations used for financial reporting, accounting analysis, and managerial decision support.	- KPI dashboards for managerial overviews. - Parallel coordinates and scatter views for ratio analysis. - Network views for intercompany flows or audit trails.
User intentions UserIntention $\sqsubseteq$ Tax4UIVB	Analytical intentions in the F&A decision cycle.	- UL_D0: periodic KPI reporting. - UL_HC: variance and ratio testing. - UL_DID_MND: capital-allocation and what-if analyses.
Visualization effects VTE $\sqsubseteq$ Tax4UIVB	Perceptual and interpretive goals used to explain taxonomy-derived score assignments.	- Visibility of accounting patterns and outliers. - Interpretability of financial reports. - Insight into dependencies, clusters, and trends.
Interaction possibilities InteractionPoss	Analyst actions performed on F&A datasets during exploration and decision support.	- Period and account selection. - Filtering by amount, category, or accounting period. - Drill-down into account hierarchies; zoom and brushing.
User groups UserGroups	Managerial and analytical roles involved in F&A reporting and decision-making.	- Top managers and CFOs: strategic decisions, regulatory reporting. - Tactical managers: departmental variance analysis. - Operative managers: process control and anomaly detection. - Analytical staff and auditors: detailed analysis.

4.7. Decision-Support Extension and Decision-Making Mechanism

The taxonomy-derived backbone described in Sections 4.2–4.5 provides the conceptual dimensions required for visualization selection, but it does not by itself define a queryable recommendation mechanism. To make the selection process operational at the ABox level, the ontology adds three decision-support classes: *DecisionContext*, *Recommendation*, and *WeightMatrixEntry*. A *DecisionContext* individual represents a concrete selection situation, a *Recommendation* individual represents one ranked candidate technique for that situation, and a *WeightMatrixEntry* individual materializes the underlying suitability value used to rank candidate techniques.

The ontology uses two complementary links between a decision context and a visualization technique. The object property *recommendsTechnique* is used as a shortcut assertion from a *DecisionContext* individual to the top-ranked technique for that context. In contrast, *Recommendation* individuals represent the full ranked candidate records: each recommendation individual links a decision context to one candidate technique and stores the associated score and rank through data properties, such as *recommendationScore* and *hasRank*. Thus, *recommendsTechnique* supports direct top-1 retrieval, whereas the *Recommendation* individuals support complete ranked-list retrieval and evaluation.

Let

$$\mathcal{T} = T_I, \quad \mathcal{U} = U_I, \quad \mathcal{G} = G_I, \quad \mathcal{J} = J_I. \quad (21)$$

These denote the ABox-level working sets of technique, user-intention, user-group, and interaction individuals introduced by the representative maps ρ_T , ρ_U , ρ_G , and ρ_J . In the current artifact these sets have cardinalities

$$|\mathcal{T}| = 13, \quad |\mathcal{U}| = 3, \quad |\mathcal{G}| = 4, \quad |\mathcal{J}| = 5.$$

The populated decision contexts are defined by a canonical-interaction assignment

$$\gamma : \mathcal{U} \times \mathcal{G} \longrightarrow \mathcal{J}, \quad (22)$$

so that each intention–user-group pair is associated with one representative interaction requirement. The context set is therefore

$$\mathcal{K} = \{ \kappa_{u,g} = (u, g, \gamma(u, g)) \mid u \in \mathcal{U}, g \in \mathcal{G} \}, \quad |\mathcal{K}| = |\mathcal{U}| \cdot |\mathcal{G}| = 12. \quad (23)$$

Although the possible interaction-aware product space contains

$$|\mathcal{U} \times \mathcal{G} \times \mathcal{J}| = 3 \times 4 \times 5 = 60$$

triples, the present release populates one canonical interaction for each intention–user-group pair. Thus, $\mathcal{K}$ is a selected 12-context subset of the larger interaction-aware space, rather than an exhaustive enumeration of all possible interaction configurations.

The function γ is a design-time assignment rather than an inferred OWL relation. For each intention–user-group pair, the canonical interaction was assigned manually from the TaxUI&BV4FADA-derived interpretation of the corresponding decision context. Thus, a populated context records one representative interaction requirement for retrieval and explanation, but the present release does not reason over all possible interaction configurations.

For each context $\kappa_{u,g} \in \mathcal{K}$, the corresponding `DecisionContext` individual is connected to its components by object-property assertions of the form

$$\begin{aligned} (\kappa_{u,g}, \text{evaluatedForCriterion}, u) &\in A_{\text{obj}}, \\ (\kappa_{u,g}, \text{evaluatedForUserGroup}, g) &\in A_{\text{obj}}, \\ (\kappa_{u,g}, \text{considersInteraction}, \gamma(u, g)) &\in A_{\text{obj}}. \end{aligned} \quad (24)$$

The shortcut top-ranked technique is represented by

$$(\kappa_{u,g}, \text{recommendsTechnique}, v^*(u, g)) \in A_{\text{obj}}, \quad (25)$$

where $v^*(u, g)$ denotes the highest-ranked technique for that context.

The suitability matrix introduced in Section 3.6 is defined over techniques, intentions, and user groups:

$$M : \mathcal{T} \times \mathcal{U} \times \mathcal{G} \longrightarrow [0, 1]. \quad (26)$$

Each value $M(v, u, g)$ is reified in the `ABox` by a `WeightMatrixEntry` individual $w_{v,u,g}$ satisfying

$$\begin{aligned} (w_{v,u,g}, \text{entryTechnique}, v) &\in A_{\text{obj}}, \\ (w_{v,u,g}, \text{entryCriterion}, u) &\in A_{\text{obj}}, \\ (w_{v,u,g}, \text{entryUserGroup}, g) &\in A_{\text{obj}}, \\ (w_{v,u,g}, \text{matrixWeight}, M(v, u, g)) &\in A_{\text{data}}. \end{aligned} \quad (27)$$

Since $|\mathcal{T}| = 13$, $|\mathcal{U}| = 3$, and $|\mathcal{G}| = 4$, the full matrix contains

$$|\mathcal{T}| \cdot |\mathcal{U}| \cdot |\mathcal{G}| = 3 \times 4 \times 13 = 156$$

weight entries.

For each technique $v \in \mathcal{T}$ and each context $\kappa_{u,g} \in \mathcal{K}$, the ABox also contains a Recommendation individual $r_{v,u,g}$ that connects the context to the candidate technique and stores its score and rank:

$$\begin{aligned}(r_{v,u,g}, \text{hasVisualizationTechnique}, v) &\in A_{\text{obj}}, \\(r_{v,u,g}, \text{evaluatedContext}, \kappa_{u,g}) &\in A_{\text{obj}}, \\(r_{v,u,g}, \text{recommendationScore}, M(v, u, g)) &\in A_{\text{data}}, \\(r_{v,u,g}, \text{hasRank}, \text{rank}_{u,g}(v)) &\in A_{\text{data}}.\end{aligned}\quad (28)$$

Thus, `WeightMatrixEntry` individuals store the underlying technique–intention–user-group suitability values, whereas `Recommendation` individuals materialize the ranked candidate records for each populated decision context.

Fix a context $(u, g) \in \mathcal{U} \times \mathcal{G}$ and let $\prec_{\mathcal{T}}$ denote a fixed total tie-breaking order on $\mathcal{T}$, for example, the lexicographic order on the technique identifiers. The ranked list is

$$R_{u,g} = \text{sort}_{\downarrow, \prec_{\mathcal{T}}}(\mathcal{T}, M(\cdot, u, g)), \quad (29)$$

i.e., the techniques are sorted by descending M -score, with ties broken by $\prec_{\mathcal{T}}$. The selected technique and the top- k recommendation set are then

$$v^*(u, g) = R_{u,g}[1], \quad \text{Top}_k(u, g) = \{R_{u,g}[1], \dots, R_{u,g}[k]\}, \quad k \in \{1, \dots, 13\}. \quad (30)$$

The rank $\text{hasRank}(r_{v,u,g})$ stored in the ABox is by construction the position of v in $R_{u,g}$. The stored `Recommendation` individuals materialize this ranked ordering in the ABox, while `recommendsTechnique` records only the first element of the ordering as a direct shortcut assertion.

The tie-breaking order is included to make the ranking function total and reproducible under future updates. If two or more techniques receive the same matrix score in a later version of the artifact, their `hasRank` values are determined by the fixed order $\prec_{\mathcal{T}}$ after sorting by descending score. In the current release, the top-ranked recommendations used in the worked scenarios are not affected by ties.

The lookup defined above specifies the intended recommendation logic. In the implementation, the relevant contexts, weight entries, scores, and ranks are stored in the ABox, and the actual $\arg \max$ and Top- k retrieval are performed by SPARQL queries or external procedures, since OWL DL itself does not compute numerical optimisation over scores.

Figure 8 summarizes the decision-support extension and highlights the distinction between the top-1 shortcut relation and the full-ranked recommendation records.

4.8. Evaluation and Validation Layer

The evaluation layer is composed of the four classes

$$\begin{aligned}C^{\text{eval}} = \{ &\text{CompetencyQuestion}, \text{ValidationQuery}, \\ &\text{EvaluationRun}, \text{RecommendationEvaluation}\}.\end{aligned}\quad (31)$$

This layer is linked to the rest of $\mathcal{O}$ by three object properties: `usesValidationQuery`, `evaluatedContext`, and `hasEvaluationResult`.

Competency questions. Let $\mathcal{Q} = \{q_1, \dots, q_n\}$ be the set of competency questions, materialized as ABox individuals $\text{CQ}_1, \dots, \text{CQ}_n$ together with their associated SPARQL queries $\text{VQ_CQ}_1, \dots, \text{VQ_CQ}_n$. The ontology stores, for each $k \in \{1, \dots, n\}$, the data triples

$$\begin{aligned}
 (CQ_k, \text{questionText}, \tau_k^q) &\in A_{\text{data}}, \\
 (CQ_k, \text{expectedAnswerText}, \tau_k^a) &\in A_{\text{data}}, \\
 (VQ_CQ_k, \text{queryText}, \tau_k^{\text{sparql}}) &\in A_{\text{data}}, \\
 (CQ_k, \text{validationPassed}, b_k) &\in A_{\text{data}},
 \end{aligned}
 \tag{32}$$

with $\tau_k^q, \tau_k^a \in \Sigma^*$, $\tau_k^{\text{sparql}} \in \Sigma_{\text{SPARQL}}^*$, and $b_k \in \mathbb{B}$. Each competency question is linked to its validation query by the ABox assertion

$$\begin{aligned}
 (CQ_k, \text{usesValidationQuery}, \\
 VQ_CQ_k) &\in A_{\text{obj}}.
 \end{aligned}
 \tag{33}$$

The current implementation has $n = 8$, hence $|I^{\text{CQ}\&\text{VQ}}| = 2n = 16$ (cf. (5)). The eight questions cover six distinct facets that together exercise the decision-support and validation layers: single-best-answer retrieval (CQ1, CQ2), top- k retrieval (CQ3), interaction-aware retrieval (CQ4), matrix-level cardinality (CQ5), threshold-based retrieval (CQ6), full-ranking retrieval (CQ7), and structural and completeness retrieval (CQ5, CQ8).

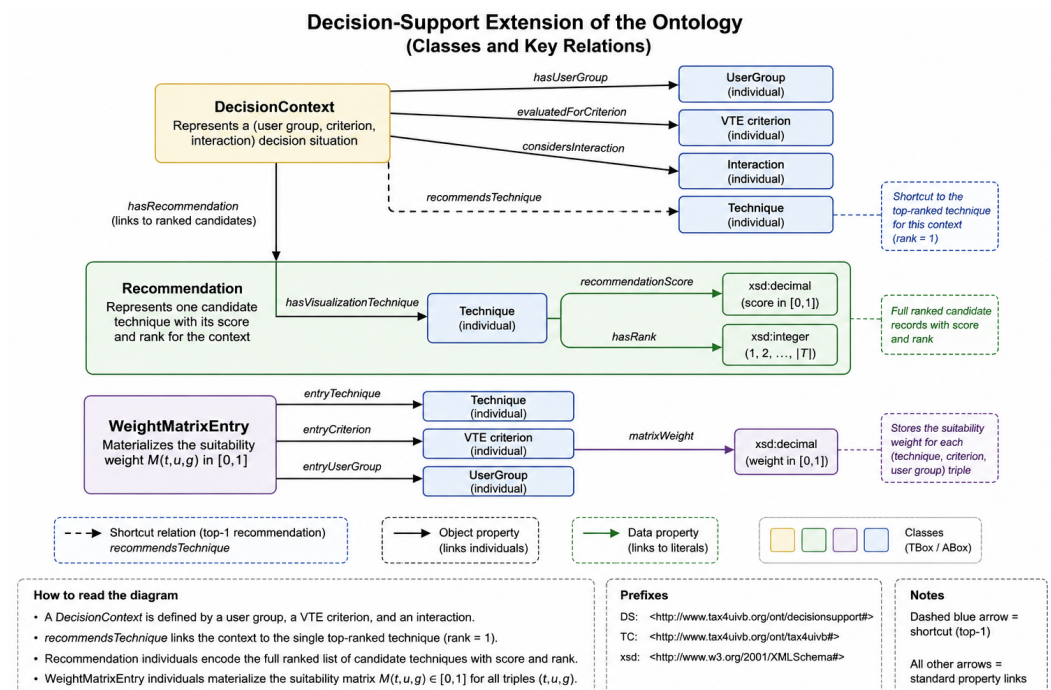


Figure 8. Decision-support extension: **DecisionContext** (user group + criterion + interaction), **Recommendation** (ranked candidates with scores and ranks), and **WeightMatrixEntry** (suitability scores). Dashed arrow: top-1 shortcut recommendsTechnique.

Note that $\text{validationPassed}(CQ_k)$ is a stored documentation of the outcome of executing the corresponding SPARQL query, not a logically derived OWL entailment. The SPARQL execution itself is performed externally, as described in Section 3.5.

Implementation-fidelity measures. For each context $(cr, ug) \in \mathcal{U} \times \mathcal{G}$, let $E_k(cr, ug) \subseteq \mathcal{T}$ denote the taxonomy-derived expected top- k set for that context, and let $t_{\text{exp}}^*(cr, ug)$ denote the expected top-1 technique. The five measures stored on the corresponding evaluation individual $e_{cr,ug} \in I^{\text{EVAL}}$ are used here as internal implementation-fidelity checks. They compare the ranked list materialized in the ABox against the taxonomy-derived expected ordering used to initialize the matrix. They should therefore not be interpreted as empirical measures of recommendation quality, predictive performance,

or user-validated effectiveness. This distinction is necessary because the expected alternatives and the matrix scores originate from the same TaxUI&BV4FADA-derived specification; the purpose of the measures is to verify correct population and retrieval of the recommendation layer.

The measures are defined as

$$P@k(cr, ug) = \frac{|\text{Top}_k(cr, ug) \cap E_k(cr, ug)|}{k}, \quad k \in \{1, 3\}, \quad (34)$$

$$MRR(cr, ug) = \frac{1}{\text{rank}(t_{\text{exp}}^*(cr, ug); R_{cr, ug})}, \quad (35)$$

$$DCG@k(cr, ug) = \sum_{i=1}^k \frac{\text{rel}_i(cr, ug)}{\log_2(i+1)}, \quad k = 3, \quad (36)$$

$$\text{nDCG}@k(cr, ug) = \frac{DCG@k(cr, ug)}{IDCG@k(cr, ug)}, \quad (37)$$

where $\text{rel}_i(cr, ug) = 1$ if the i th-ranked technique in $R_{cr, ug}$ belongs to $E_k(cr, ug)$ and 0 otherwise, and $IDCG@k$ is the value of $DCG@k$ on the ideal ranking. The aggregate coverage is

$$\text{Coverage} = \frac{|\{(cr, ug) \in \mathcal{U} \times \mathcal{G} : \text{Top}_1(cr, ug) \neq \emptyset\}|}{|\mathcal{U} \times \mathcal{G}|}. \quad (38)$$

In the current full-matrix release, Coverage = 1 is expected because every one of the twelve intention–user-group contexts is populated with thirteen candidate recommendations. The measure is therefore retained as a regression and completeness check for future partial, extended, or incrementally updated ABox versions, rather than as an independent indicator of recommendation quality.

Recommendation-evaluation tuple. Each evaluation individual $e \in I^{\text{EVAL}}$ stores

$$\Phi(e) = (P@1, P@3, MRR, \text{nDCG}@3, \text{cov}, \tau_{\text{cons}})(e) \in [0, 1]^6, \quad (39)$$

together with $\text{qualityLabel}(e) \in \{\text{LOW}, \text{MODERATE}, \text{HIGH}\}$ and $\text{evaluationNote}(e) \in \Sigma^*$. Each evaluation is linked to a context via $(e, \text{evaluatedContext}, \kappa) \in A_{\text{obj}}, \kappa \in I^{\text{CTX}}$, and gathered under a single run via $(e, \text{hasEvaluationResult}, \text{EvalRun_FullMatrix})$, yielding $|I^{\text{EVAL}}| = |I^{\text{CTX}}| = 12$. The sixth component τ_{cons} is stored in the OWL file under the data property `taxonomyConsistencyScore` and is interpreted as a taxonomy-consistency score rather than as an independent expert-validation outcome (no separate expert elicitation was performed). In the current heuristic full-matrix initialization, $P@1 = P@3 = MRR = \text{nDCG}@3 = \text{cov} = 1.0$ uniformly across all twelve contexts and $\tau_{\text{cons}} \in [0.90, 0.98]$, with $\text{qualityLabel} = \text{HIGH}$ uniformly; these values indicate implementation fidelity rather than empirical accuracy (see Section 5.4).

4.9. Example Recommendation Scenarios

Two contrasting F&A scenarios are instantiated as ABox triples $(cr, ug, \gamma(cr, ug))$ and resolved through the selection function (30). The two scenarios share the same ontology and the same decision-support pattern, yet yield different top-ranked techniques. This is a non-degeneracy property of the recommendation function v^* : there exist contexts at which the top-ranked techniques differ. We do not claim that v^* is injective on $\mathcal{U} \times \mathcal{G}$, and indeed, the populated matrix assigns the same top-ranked technique to several contexts.

4.9.1. Scenario 1: Financial Analyst Testing a Gross-Margin Hypothesis

Maria is a financial analyst (mapped to `UG_AnalyticalStaff` via σ in Table 5). Her decision context is

$$\kappa_1 : (cr, ug, \gamma(cr, ug)) = (Crit_UL_HC, UG_AnalyticalStaff, Int_SelectDataAttributes), \quad (40)$$

realized as $\kappa_1 = CTX_UL_HC_AnalyticalStaff$. The selection function returns

$$\begin{aligned} v^*(Crit_UL_HC, UG_AnalyticalStaff) &= Tech_ParCor, \\ M(Tech_ParCor, Crit_UL_HC, UG_AnalyticalStaff) &= 0.98, \end{aligned} \quad (41)$$

materialized as $r_1 = REC_UL_HC_AnalyticalStaff_ParCor$ with $hasRank(r_1) = 1$.

Table 5. Instance-level decomposition for Scenario 1 (financial analyst testing a gross-margin hypothesis), as encoded in the OWL artifact *FinAccOnt.ver-01*.

Modeling Element	OWL Class/Property	Concrete Instance/Value
User-group assignment	UserGroups/ evaluatedForUserGroup	UG_AnalyticalStaff
Decision context	DecisionContext	CTX_UL_HC_AnalyticalStaff
Analytical-intention link	evaluatedForCriterion	→ Crit_UL_HC
User-group link	evaluatedForUserGroup	→ UG_AnalyticalStaff
Interaction link	considersInteraction	→ Int_SelectDataAttributes
Top recommended technique	recommendsTechnique	→ Tech_ParCor
Recommendation individual	Recommendation/ hasContext	REC_UL_HC_AnalyticalStaff_ParCor
Weight-matrix entry	WeightMatrixEntry/ hasWeightEntry	WME_UL_HC_AnalyticalStaff_ParCor
Matrix weight	matrixWeight	0.98
Recommendation score	hasWeight/ recommendationScore	0.98
Recommendation rank	hasRank	1
Context score	hasScore	0.98
Quantitative evaluation	RecommendationEvaluation/ evaluatedContext	EVAL_UL_HC_AnalyticalStaff (P@1 = P@3 = MRR = nDCG@3 = Coverage = 1.0; $\tau_{cons} = 0.97$; label = high).

The full ranked set $Top_5(Crit_UL_HC, UG_AnalyticalStaff)$ is reported in Table 6.

Table 6. Top-ranked techniques for Scenario 1 ($CTX_UL_HC_AnalyticalStaff$), derived from the $REC_UL_HC_AnalyticalStaff_*$ individuals in the OWL file.

Rank	Recommended Technique	Score
1	Tech_ParCor (parallel-coordinates view of financial ratios)	0.98
2	Tech_HierarParCor (hierarchical parallel-coordinates)	0.91
3	Tech_Scatterplots (cross-account scatter analysis)	0.91
4	Tech_HierarScatterplots	0.87
5	Tech_Pixel_oriented_display (transaction-level pixel display)	0.86

4.9.2. Scenario 2: CFO Seeking a Quarterly Data Overview

The CFO (Tables 7 and 8) is mapped to $UG_TopManagers$ and $Crit_UL_D0$ through σ , and the decision context is

$$\kappa_2 : (cr, ug, \gamma(cr, ug)) = (Crit_UL_D0, UG_TopManagers, Int_SelectSubset), \quad (42)$$

realized as $\kappa_2 = CTX_UL_D0_TopManagers$. The selection function returns

$$\begin{aligned} v^*(\text{Crit_UL_D0}, \text{UG_TopManagers}) &= \text{Tech_Dashboards}, \\ M(\text{Tech_Dashboards}, \text{Crit_UL_D0}, \text{UG_TopManagers}) &= 0.99, \end{aligned} \quad (43)$$

materialized as $r_2 = \text{REC_UL_D0_TopManagers_Dashboards}$ with $\text{hasRank}(r_2) = 1$.

Table 7. Instance-level decomposition for Scenario 2 (CFO seeking a quarterly data overview), as encoded in the OWL artifact `FinAccOnt.ver-01`.

Modeling Element	OWL Class/Property	Concrete Instance/Value
User-group assignment	UserGroups/ evaluatedForUserGroup	UG_TopManagers
Decision context	DecisionContext	CTX_UL_D0_TopManagers
Analytical-intention link	evaluatedForCriterion	$\rightarrow \text{Crit_UL_D0}$
User-group link	evaluatedForUserGroup	$\rightarrow \text{UG_TopManagers}$
Interaction link	considersInteraction	$\rightarrow \text{Int_SelectSubset}$
Top recommended technique	recommendsTechnique	$\rightarrow \text{Tech_Dashboards}$
Recommendation individual	Recommendation/ hasContext	REC_UL_D0_TopManagers_Dashboards
Weight-matrix entry	WeightMatrixEntry/ hasWeightEntry	WME_UL_D0_TopManagers_Dashboards
Matrix weight	matrixWeight	0.99
Recommendation score	hasWeight/ recommendationScore	0.99
Recommendation rank	hasRank	1
Context score	hasScore	0.99
Quantitative evaluation	RecommendationEvaluation/ evaluatedContext	EVAL_UL_D0_TopManagers ($P@1 = P@3 = \text{MRR} = \text{nDCG}@3 = \text{Coverage} = 1.0$; $\tau_{\text{cons}} = 0.90$; label = high).

Table 8. Top-ranked techniques for Scenario 2 (CTX_UL_D0_TopManagers), derived from the REC_UL_D0_TopManagers_* individuals in the OWL file.

Rank	Recommended Technique	Score
1	Tech_Dashboards (financial KPI dashboard)	0.99
2	Tech_Glyphs (glyph-based portfolio summaries)	0.77
3	Tech_HierarGlyphs (hierarchical glyph summaries)	0.72
4	Tech_Scatterplots (cross-account scatter analysis)	0.70
5	Tech_Polar_coordinates (cyclic/seasonal display)	0.66

4.9.3. Non-Degeneracy of the Recommendation Function

Both scenarios share the same TBox vocabulary ($C^{\text{tax}} \cup C^{\text{dec}}$), the same decision-support pattern (24)–(28), and the same SPARQL query template. Yet there exist contexts $(cr_1, ug_1), (cr_2, ug_2) \in \mathcal{U} \times \mathcal{G}$ with

$$v^*(cr_1, ug_1) \neq v^*(cr_2, ug_2), \quad (44)$$

witnessed concretely by the two worked scenarios:

$$\begin{aligned} v^*(\text{Crit_UL_HC}, \text{UG_AnalyticalStaff}) &= \text{Tech_ParCor}, \\ v^*(\text{Crit_UL_D0}, \text{UG_TopManagers}) &= \text{Tech_Dashboards}. \end{aligned}$$

Moreover, the two top-ranked techniques do not appear in each other's top-five sets:

$$\begin{aligned}\text{Tech_ParCor} &\notin \text{Top}_5(\text{Crit_UL_D0}, \text{UG_TopManagers}), \\ \text{Tech_Dashboards} &\notin \text{Top}_5(\text{Crit_UL_HC}, \text{UG_AnalyticalStaff}).\end{aligned}$$

These observations establish that the recommendation function v^* is not globally constant: the populated weight matrix encodes meaningful F&A discrimination across at least the two contrasting contexts shown. They do not entail that v^* is injective on the full 3×4 context grid, and the populated matrix indeed assigns the same top-ranked technique to several contexts.

4.10. Upper Ontology Alignment

To support future interoperability, an explanatory mapping

$$\mu : C^{\text{tax}} \cup C^{\text{dec}} \cup C^{\text{eval}} \longrightarrow \mathcal{B} \times \{\text{Cont}, \text{Occ}\} \quad (45)$$

is provided, where $\mathcal{B} = \{\text{ICE}, \text{PP}, \text{EP}, \dots\}$ is a fixed finite set of upper-ontology category labels drawn from the Basic Formal Ontology (BFO), the Information Artifact Ontology (IAO), and the Ontology for Biomedical Investigations (OBI), and $\{\text{Cont}, \text{Occ}\}$ distinguishes continuants from occurrents (Figure 9). Each label in $\mathcal{B}$ is a string identifier of an external upper-ontology class, not an OWL class in $\mathcal{O}$; the map μ is therefore a meta-level annotation function and does not change the OWL entailment regime of the current artifact. In particular, μ is not realized in the OWL file through `owl:imports`, `owl:equivalentClass`, or formal subclass axioms to BFO/IAO/OBI entities.

Within the continuant branch, BFO further partitions entities into independent, specifically dependent, and generically dependent continuants. The four taxonomy-derived classes C^{tax} are interpreted as information-content entities, $\mu(c) = (\text{ICE}_{\text{IAO}}, \text{Cont}_{\text{GDC}})$ for every $c \in C^{\text{tax}}$. The decision-support classes are also interpreted as information-content entities: a `DecisionContext` records a structured decision situation, a `Recommendation` records one ranked candidate technique, and a `WeightMatrixEntry` records one suitability value in the taxonomy-derived matrix. The numerical value stored through `matrixWeight` denotes a suitability score, but the `WeightMatrixEntry` individual itself is treated as an information-content record rather than as a specifically dependent quality. The dynamic part of the evaluation is represented by `EvaluationRun`, which is aligned with a planned process in OBI. By contrast, `RecommendationEvaluation` is treated as an information-content entity that records the outcome of such an evaluation run, including metric values, such as `precisionAt1`, `ndcgAt3`, `taxonomyConsistencyScore`, and `qualityLabel`. Table 9 enumerates μ explicitly.

The current OWL file does not contain a class `VisualizationSelectionProcess`; visualization selection is therefore represented indirectly through `DecisionContext`, `Recommendation`, `WeightMatrixEntry`, and the evaluation classes. The process-oriented reading is supplied primarily by `EvaluationRun`, which represents the execution of the validation and evaluation procedure. The resulting metric values are then recorded on `RecommendationEvaluation` individuals, which are interpreted as evaluation-result records rather than as processes themselves. `DecisionContext` is likewise modeled as an information-content entity: a structured aggregation of decision parameters rather than a temporal entity.

Table 9. Conceptual alignment of ontology classes with upper-ontology categories. The alignment is interpretive: no owl:imports, owl:equivalentClass, or formal subclass axioms to BFO, IAO, or OBI entities are asserted in the current OWL file.

Ontology Class	Upper-Ontology Reading	Source
Continuant/information-content records		
MDMV_techniques	Information Content Entity/GDC	IAO/BFO
Tax4UIVB	Information Content Entity/GDC	IAO/BFO
InteractionPoss	Information Content Entity/GDC	IAO/BFO
UserGroups	Information Content Entity/GDC	IAO/BFO
DecisionContext	Information Content Entity/GDC	IAO/BFO
Recommendation	Information Content Entity/GDC	IAO/BFO
WeightMatrixEntry	Information Content Entity/GDC	IAO/BFO
CompetencyQuestion	Information Content Entity/GDC	IAO/BFO
ValidationQuery	Information Content Entity/GDC	IAO/BFO
RecommendationEvaluation	Information Content Entity/GDC	IAO/BFO
Occurrent/evaluation process		
EvaluationRun	Planned Process/Occurrent	OBI/BFO

Per-class subtleties. *WeightMatrixEntry* is treated as an information-content record even though the value stored on *matrixWeight* denotes a suitability score; the individual is a record, not the quality itself. *RecommendationEvaluation* records the outcome of an *EvaluationRun* (metric values, τ_{cons} , *qualityLabel*). All other classes are conceptual categories for their respective taxonomy dimensions.

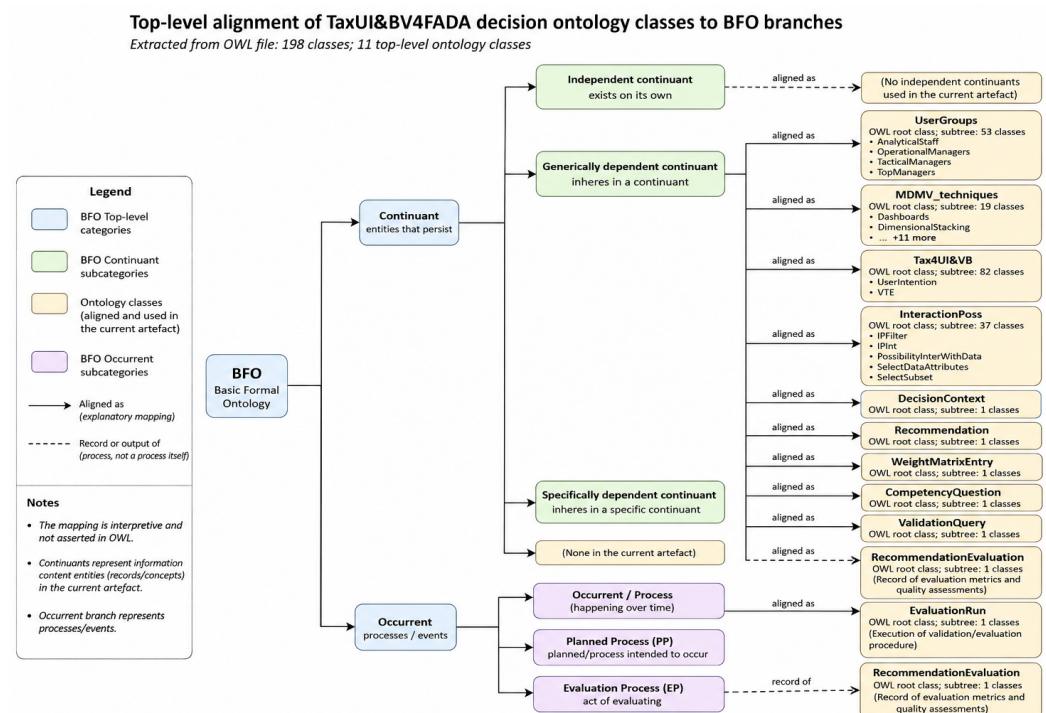


Figure 9. Top-level placement of selected ontology classes in the BFO/IAO/OBI hierarchy (continuant vs. occurrent split; see Section 4.10).

4.11. Status of the Framework and Open Items

The artifact presented here is an ontology-based decision-support framework, not a deployed end-to-end visualization recommender. Formally, what is fully populated is the 9-tuple

$$\mathcal{F} = \langle \mathcal{T}, \mathcal{U}, \mathcal{G}, \mathcal{J}, \mathcal{K}, M, R, \mathcal{Q}, \mathcal{E} \rangle, \quad (46)$$

where $\mathcal{T}, \mathcal{U}, \mathcal{G}, \mathcal{J}$ are the working sets defined in (21) with cardinalities 13, 3, 4, 5; $\mathcal{K}$ is the set of $|\mathcal{K}| = 12$ decision contexts defined in (23); M is the populated suitability matrix in (26) with 156 entries in $[0.40, 0.99]$; $R = \{R_{cr,ug}\}_{(cr,ug) \in \mathcal{U} \times \mathcal{G}}$ is the ranking-function family

in (29), realized at the ABox as $|I^{\text{REC}}| = 156$ ranked recommendation individuals; $\mathcal{Q}$ is the set of $|\mathcal{Q}| = 8$ competency questions with their SPARQL validation queries, covering six retrieval and structural facets (Section 4.8); and $\mathcal{E}$ is the set of $|\mathcal{E}| = |I^{\text{EVAL}}| = 12$ internal-consistency evaluations gathered under EvalRun_FullMatrix.

The context component $\mathcal{K}$ should be interpreted carefully. Although the possible interaction-aware space contains

$$|\mathcal{U} \times \mathcal{G} \times \mathcal{J}| = 3 \times 4 \times 5 = 60$$

combinations, the current release populates only

$$\mathcal{K} \subset \mathcal{U} \times \mathcal{G} \times \mathcal{J}, \quad |\mathcal{K}| = 12,$$

with one canonical interaction assigned to each intention–user–group pair. Thus, the artifact supports interaction-aware interpretation through canonical decision contexts, but it does not yet exhaustively populate or reason over all possible interaction configurations.

The visualization-effect branch VTE is also explanatory in the present release and is not an independent dimension of the suitability matrix. The current matrix remains

$$M : \mathcal{T} \times \mathcal{U} \times \mathcal{G} \rightarrow [0, 1],$$

while VTE classes provide vocabulary for interpreting the taxonomy-derived rationale behind score assignments. A future extension could define an enriched matrix that conditions scores on visualization effects, but this is outside the scope of the current artifact.

Several items remain open and define the scope of the next implementation step. The matrix entries $M(v, cr, ug)$ are currently produced by the heuristic, taxonomy-consistent initialization described in Section 3.6, and have not been validated against an independent expert elicitation. The data property `taxonomyConsistencyScore` records the taxonomy-consistency score τ_{cons} and takes values in $[0.90, 0.98]$; its semantics in the present release is internal alignment with the taxonomy rather than independent expert judgement. Replacing the heuristic weights with expert-elicited or empirically estimated ones is therefore the most direct route to strengthening the recommendation layer.

Rule-based inference in the current artifact is limited to three SWRL rules (twelve SWRL atoms in total) that fire the `recommendsTechnique` link for three anchor decision contexts: $(\text{Crit_UL_DID_MND}, \text{UG_TacticalManagers}) \rightarrow \text{Tech_DimensionalStacking}$, $(\text{Crit_UL_D0}, \text{UG_TopManagers}) \rightarrow \text{Tech_Dashboards}$, and $(\text{Crit_UL_HC}, \text{UG_AnalyticalStaff}) \rightarrow \text{Tech_ParCor}$. These rules provide limited anchor assertions only. They do not compute the principal ranking step: the $\arg \max$ operation, the top- k retrieval, and the full ranked lists are obtained from the stored ABox scores and ranks by SPARQL queries or external procedures, not by OWL DL or SWRL reasoning.

Consequently, the current artifact should be read as a semantic foundation for recommendation retrieval rather than as a complete recommendation system. It provides the ontology vocabulary, populated ABox, suitability matrix, ranked recommendation records, competency questions, validation queries, and internal implementation-fidelity checks. It does not yet include a live F&A data-ingestion pipeline, a user-facing recommender interface, empirical validation with financial and accounting users, SHACL-based closed-world constraint checking, or expert-elicited weighting. These items define the immediate future work path for moving from the present ontology framework toward a deployed F&A visualization recommender.

5. Internal Validation and Consistency Analysis

The proposed ontology is evaluated as a single artifact using established ontology-engineering practices. Throughout this section, the artifact is understood as the 10-tuple O defined in (1), with class set C , individual set I , property sets P_o, P_d , and the assertion and rule sets $A_{\text{sub}}, A_{\text{type}}, A_{\text{obj}}, A_{\text{data}}, A_{\text{ann}}, R_{\text{swrl}}$. The evaluation focuses on four complementary dimensions — structural metrics, logical consistency, competency-question validation, and recommendation-matrix consistency—and is grounded strictly in elements explicitly present in the OWL implementation; no claims are made about functionality that is not encoded in the file. A closing subsection states what the evaluation does and does not establish.

5.1. Structural Metrics

The ontology is organized around eleven top-level named classes, partitioned into the four taxonomy-derived classes of C^{tax} —which encode the principal conceptual dimensions of the domain (visualization techniques, classification criteria, interaction possibilities, and user groups)—the three decision-support classes of C^{dec} , and the four evaluation classes of C^{eval} (Section 4.1). The OWL/RDF implementation contains $|C| = 198$ classes connected through $|A_{\text{sub}}| = 187$ subclass axioms, so the pair (C, A_{sub}) forms a finite, acyclic, rooted partial order of depth $\max_{c \in C} d(c) = 5$. The structure is well-formed: no cycles and no multiple-inheritance axioms are present. A detailed overview is given in Table 10.

Table 10. Structural metrics of the submitted OWL/RDF artifact. Row groupings correspond to the decision-support layer (Section 4.7), the evaluation layer (Section 4.8), and the SWRL layer (Section 4.11).

Metric	Observed Value
<i>Class hierarchy (TBox)</i>	
Total number of classes	198
Number of top-level classes	11
of which taxonomy-derived	4
of which decision-support	3
of which evaluation-related	4
Number of subclass axioms	187
Maximum hierarchy depth	5
Number of leaf classes	149
Number of internal classes	49
Average branching factor	3.82
Maximum branching factor	13 under MDMV_techniques
Classes with multiple inheritance	0
Equivalent-class axioms	0
Disjoint-class axioms	0
Restriction expressions	0
<i>Properties</i>	
Declared object properties	18
Declared data properties	19
Declared annotation properties	4
Inverse-property axioms	0
<i>Individuals (ABox)</i>	
Declared named individuals	390
Tech_* technique individuals	13
Crit_* criterion/intention individuals	3
UG_* user-group individuals	4
Int_* interaction individuals	5
CTX_* decision-context individuals	12
WME_* weight-matrix entries	156
REC_* ranked recommendation individuals, full matrix	156
EVAL_* recommendation-evaluation individuals	12
CQ* competency-question individuals	8
VQ_CQ* validation-query individuals	8
EvalRun_FullMatrix	1
legacy DC_*/REC_* individuals	4
punned class-individual pairs	8

Table 10. Cont.

Metric	Observed Value
<i>Axiom counts (ABox and annotations)</i>	
Class assertions (<code>rdf:type</code>)	404
Object-property assertions	1655
Data-property assertions	1102
Annotation assertions	675
<i>Rule layer</i>	
SWRL rules (<code>swrl:Imp</code>)	3
SWRL atoms, class and property predicates	12

Note. The structural metrics describe the submitted OWL artifact and should not be read as empirical validation of the recommendation model.

The ontology exhibits a predominantly hierarchical structure in which most classes are leaves, indicating fine-grained specialization of concepts. The moderate hierarchy depth and balanced branching factor suggest a well-controlled conceptual decomposition. Two structural characteristics deserve explicit comment, because they bound the formal expressiveness of the artifact. First, the TBox contains zero `owl:disjointWith/AllDisjointClasses` axioms, zero `owl:equivalentClass` axioms, and zero `owl:Restriction` expressions, and the property set contains zero `owl:inverseOf` axioms. Logically, this positions the artifact in the lightweight RDFS-plus/OWL RL fragment of OWL DL: it supports class-hierarchy reasoning and ABox typing, but it does not encode constraint-style axioms that would let an OWL DL reasoner derive contradictions or new class memberships from conjunctive or restriction-based patterns. Second, the rule layer R_{SWRL} consists of three SWRL rules with twelve atoms in total (Section 4.11); these fire the `recommendsTechnique` link for three anchor decision contexts and do not compute the principal arg max or top-*k* selection, which is performed externally by SPARQL queries. Consequently, the present artifact behaves as a rich, queryable vocabulary plus rule-anchored assertions rather than as a constraint-rich DL ontology, and this is the appropriate reading of any claim of “OWL DL reasoning” made elsewhere in this paper.

5.2. Reasoner-Based Structural Consistency Check

Logical consistency was verified with the HermiT reasoner integrated in Protégé Desktop, using default settings, as described in Section 3.5. The reasoner reported no inconsistencies and no unsatisfiable classes for the fully populated artifact; class-hierarchy classification and ABox realization completed without errors. Given the structural profile reported in Table 10—in particular the absence of disjointness, equivalence, and restriction axioms—the consistency check primarily certifies that the subsumption hierarchy and the ABox typings are syntactically and structurally well-formed, rather than excluding non-trivial logical contradictions that could not arise in this expressive fragment.

5.3. Competency-Question Validation

The query-level functionality of the populated ABox is validated through eight competency questions CQ1–CQ8, each encoded as a named individual in the OWL file. Every `CompetencyQuestion` individual stores its natural-language formulation through the data property `questionText`, its expected answer through `expectedAnswerText`, and a Boolean validation outcome through `validationPassed`. It is linked through `usesValidationQuery` to a corresponding `ValidationQuery` individual VQ_CQ1–VQ_CQ8, whose `queryText` field carries the SPARQL query that operationalizes the question. The SPARQL queries are executed externally against the RDF graph (Section 3.5); the `validationPassed` flag stored in the OWL file is documentation of the most recent execution outcome and is not an OWL DL entailment. The eight queries are reproduced in full in

Appendix B, grouped by the six facets enumerated in Section 3.3. Table 11 summarizes the encoded competency questions and their stored outcomes.

All eight `validationPassed` flags evaluate to `true`, so the SPARQL ranking layer returns the techniques specified by the design across the six facets enumerated in Section 3.3: single-best-answer (CQ1, CQ2), top-*k* (CQ3), interaction-aware (CQ4), threshold-based (CQ6), full-ranking (CQ7), and structural/completeness (CQ5, CQ8).

Table 11. Competency questions encoded in the ontology and their validation outcomes, as stored in the OWL file. Each row corresponds to a `CompetencyQuestion` individual paired with a `SPARQL ValidationQuery`; all eight queries return the expected answer when executed over the populated ABox.

ID	Question	Expected Answer	Passed
CQ1	Which visualization technique is recommended for top managers seeking data overview?	Tech_Dashboards	true
CQ2	Which technique receives the highest score for analytical staff performing hypothesis confirmation?	Tech_ParCor	true
CQ3	What are the top three techniques for tactical managers engaged in decision-making?	Tech_DimensionalStacking, Tech_HierarDimenStacking, Tech_Nodes_and_Links	true
CQ4	Which interaction mode is considered in the context of analytical staff performing hypothesis confirmation?	Int_SelectDataAttributes	true
CQ5	How many weight-matrix entries are encoded in the full recommendation matrix?	156	true
CQ6	Which techniques achieve a recommendation score of at least 0.85 for analytical staff performing hypothesis confirmation?	Tech_ParCor, Tech_HierarParCor, Tech_Scatterplots, Tech_HierarScatterplots, Tech_Pixel_oriented_display	true
CQ7	Retrieve the complete ranking, with scores and ranks, for the CFO data-overview context <code>CTX_UL_D0_TopManagers</code> .	13 ranked rows, from Tech_Dashboards (rank 1, score 0.99) down to the lowest-ranked technique (rank 13)	true
CQ8	Does every full-matrix decision context have exactly 13 candidate recommendation individuals?	empty result set (each of the 12 <code>CTX_*</code> contexts has exactly 13 recommendations; total 156)	true

5.4. Implementation-Fidelity Check of the Recommendation Matrix

In addition to competency-question validation, the ontology contains a dedicated internal-consistency layer realized through the class `RecommendationEvaluation`. The OWL file encodes twelve `EVAL_*` individuals, one per (intention, user-group) decision context. Every evaluation individual *e* is linked to its context κ and gathered under a single evaluation run by two object-property assertions:

$$(e, \text{evaluatedContext}, \kappa) \in A_{\text{obj}},$$

$$(e, \text{hasEvaluationResult}, \text{EvalRun_FullMatrix}) \in A_{\text{obj}}.$$

Each evaluation individual stores the ranking-quality measures defined formally in (34)–(38) of Section 4.8—`P@1`, `P@3`, `MRR`, `nDCG@3`, and `Coverage`—together with the taxonomy-consistency score τ_{cons} (stored under the legacy property name `expertAgreement`), a categorical `qualityLabel`, and a free-text `evaluationNote`.

Scope and intent of this check. The values reported below are internal-consistency measures of the implementation against its own specification, not an external evaluation of recommendation quality. The expected top techniques used to score each context are

the same techniques whose suitability was assigned the highest values by the heuristic, taxonomy-consistent initialization of M described in Section 3.6. A score of 1.0 on P@1, P@3, MRR, nDCG@3, or Coverage therefore certifies that the populated ABox faithfully realizes the design specification—i.e., that the SPARQL ranking layer returns the techniques the taxonomy authors intended for each decision context—not that those techniques have been independently shown to be the right recommendations for real F&A analysts. The taxonomy-consistency score τ_{cons} similarly measures internal alignment of the populated weights with the taxonomy’s expected ordering; no independent expert elicitation was performed in the present study. The complete set of stored values is given in Table 12.

Table 12. Internal consistency of the recommendation layer for the twelve (intention, user-group) decision contexts, as encoded in EVAL_* individuals collected under EvalRun_FullMatrix.

Decision Context (Intention, User Group)	P@1	P@3	MRR	nDCG@3	Cov.	τ_{cons}	Label
CTX_UL_DO_TopManagers	1.00	1.00	1.00	1.00	1.00	0.90	high
CTX_UL_DO_TacticalManagers	1.00	1.00	1.00	1.00	1.00	0.92	high
CTX_UL_DO_OperativeManagers	1.00	1.00	1.00	1.00	1.00	0.94	high
CTX_UL_DO_AnalyticalStaff	1.00	1.00	1.00	1.00	1.00	0.96	high
CTX_UL_HC_TopManagers	1.00	1.00	1.00	1.00	1.00	0.91	high
CTX_UL_HC_TacticalManagers	1.00	1.00	1.00	1.00	1.00	0.93	high
CTX_UL_HC_OperativeManagers	1.00	1.00	1.00	1.00	1.00	0.95	high
CTX_UL_HC_AnalyticalStaff	1.00	1.00	1.00	1.00	1.00	0.97	high
CTX_UL_DID_MND_TopManagers	1.00	1.00	1.00	1.00	1.00	0.92	high
CTX_UL_DID_MND_TacticalManagers	1.00	1.00	1.00	1.00	1.00	0.94	high
CTX_UL_DID_MND_OperativeManagers	1.00	1.00	1.00	1.00	1.00	0.96	high
CTX_UL_DID_MND_AnalyticalStaff	1.00	1.00	1.00	1.00	1.00	0.98	high
Mean (12 contexts)	1.00	1.00	1.00	1.00	1.00	0.94	high

Three observations follow from Table 12. First, the four ranking-quality measures (P@1, P@3, MRR, nDCG@3) and the Coverage score reach the value 1.0 uniformly across all twelve contexts. Because the expected top techniques used to score each context were taken from the TaxUI&BV4FADA taxonomy and the heuristic initialization of M assigns the highest weights to those same techniques, the all-1.0 result is a unit-test outcome: it certifies that the implemented ABox does what the design specifies, not that the design itself predicts good recommendations. Second, the τ_{cons} values lie in $[0.90, 0.98]$ with mean 0.94 and are internal to the artifact: they measure how closely the populated weight matrix follows the taxonomy’s expected ordering, with the residual gap of up to 0.10 concentrated in the TopManagers contexts. Third, the categorical qualityLabel resolves to high for every context; this is a uniform readiness label for the current release rather than a graded empirical verdict, and finer-grained labels (e.g., very-high, moderate) become meaningful only once empirically elicited rankings are integrated.

The competency-question and recommendation-matrix layers together establish that the ontology is structurally and operationally sound: SPARQL queries against I^{REC} return the techniques specified by the taxonomy, and the ABox satisfies the expected cardinality, top-1, top- k , and interaction-aware properties (CQ1–CQ8). Both layers are encoded inside the OWL file as first-class individuals, so adding new contexts or revising weights requires only updating the corresponding WeightMatrixEntry, Recommendation, and EVAL_* individuals; the validation tables in this section are then regenerated by re-running the SPARQL queries stored in VQ_CQ1–VQ_CQ8.

5.5. Limits of the Validation

The evaluation reported in Sections 5.1–5.4 establishes the internal consistency and queryability of the submitted artifact. The structural metrics confirm the intended class, property, and individual cardinalities; the reasoner-based consistency check confirms that the artifact is logically consistent within its current expressive fragment; the eight competency-question queries return their stored expected answers; and the recommendation-matrix evaluation shows that the populated ABox realizes the taxonomy-derived expected rankings across all twelve decision contexts. These results establish structural consistency, ABox retrieval functionality, and implementation fidelity of the recommendation layer.

However, this validation does not establish empirical recommendation accuracy or practical optimality for financial and accounting professionals. The current scores are taxonomy-consistency values derived from the TaxUI&BV4FADA ordering, not values learned from usage data or obtained through independent expert elicitation. Predictive validity would require an external study with practicing financial analysts, accountants, controllers, and CFOs, in which the recommended techniques are compared against expert judgments, user preferences, or task-performance outcomes. Such validation is outside the scope of the present release and is treated as a direct next step.

Four limitations follow directly from the current validation design. First, the suitability weights have not yet been independently elicited or empirically estimated. The data property `taxonomyConsistencyScore` records the taxonomy-consistency score τ_{cons} , which is not evidence of independent expert agreement. Second, no user study has yet been conducted with F&A professionals. Third, the current OWL artifact does not include SHACL constraints for closed-world validation of cardinalities, datatypes, completeness conditions, or matrix-population rules. Fourth, the financial and accounting semantics are documented through the external mapping σ rather than encoded through a separately imported financial/accounting ontology.

The artifact also imposes semantic-coverage limits that bound the recommendation function from the data side. Visualization techniques, user intentions, interaction possibilities, and user groups are explicitly represented through dedicated branches of the class hierarchy (Sections 4.2–4.5). Data types, however, are not modeled as independent ontology concepts, and no classes or properties are defined for the preprocessing, transformation, encoding, rendering, or deployment stages of a complete visualization pipeline. The artifact, therefore, covers the selection step of the visualization process—which technique to use for which intention, user-group, and canonical-interaction context—but not the upstream and downstream stages surrounding that selection. This scope is consistent with the positioning in Section 2: the present framework addresses the user-centered visualization-selection gap, not the full visualization-pipeline gap.

A further practical limitation concerns maintenance. Because visualization techniques and interaction paradigms continue to evolve, the class hierarchy and the populated weight matrix require periodic curation: each newly introduced technique entails re-scoring across the affected decision contexts, and deprecated techniques must be retired without breaking existing recommendation individuals. Keeping this maintenance effort tractable is a design concern for any deployed version of the framework, and is one motivation for the SHACL-based population checks and lightweight governance procedures discussed in Sections 6 and 8.

This limitation is mathematically tractable. Extending the model with data-characteristic constraints would amount to extending the current suitability matrix from

$$M : \mathcal{T} \times \mathcal{U} \times \mathcal{G} \longrightarrow [0,1] \quad (47)$$

to

$$M' : \mathcal{T} \times \mathcal{U} \times \mathcal{G} \times \mathcal{D} \longrightarrow [0, 1], \quad (48)$$

where $\mathcal{D}$ collects data-type and data-condition constraints, such as dimensionality, cardinality, measurement scale, distribution class, temporal structure, hierarchical structure, missingness, or uncertainty. The ranking $R_{cr,ug}$ defined in (29) would then be replaced by a data-conditioned ranking $R_{cr,ug,d}$ for each $d \in \mathcal{D}$, while the same SPARQL-retrieval principle could be retained. The required ABox extension would introduce additional individuals, such as `DataContext` instances and additional `WME_*` records per data condition. Such an extension would be compatible with the present framework, but it would require additional ontology modeling and validation and is therefore left for future work.

A further limitation concerns the financial and accounting semantic layer. In the present release, F&A semantics are documented through the external mapping σ rather than encoded as a separately imported domain ontology. This means that the ontology formalizes the visualization-selection structure for F&A decision contexts, while deeper domain reasoning over financial statements, accounting standards, XBRL facts, reporting rules, or audit-specific concepts is outside the scope of the current artifact. Future work should investigate explicit links to domain vocabularies such as XBRL, financial-reporting ontologies, or accounting-standard vocabularies.

Finally, the validation should be interpreted in light of the artifact's expressive profile. The current implementation operates as a vocabulary-plus-assertions ontology with a populated decision-support ABox, limited anchor SWRL rules, and external SPARQL-based retrieval. It is not a constraint-rich DL ontology and does not use OWL reasoning to compute numerical rankings or recommendation optimality. Thus, the validation establishes structural consistency, ABox queryability, and implementation fidelity of the populated recommendation layer, but it does not establish empirical recommendation accuracy, user-perceived usefulness, or full-pipeline coverage.

6. Discussion

This section situates the artifact with respect to its intended scope, its validation strategy, and its practical applicability. The aim is to make explicit what the framework does and does not yet claim, and to outline concretely how the internal, artifact-centered evidence reported in Section 5 could be complemented by external, user-centered validation.

6.1. Scope of the Contribution

The contribution of this work is the formalization of visualization-selection knowledge for financial and accounting (F&A) contexts as a machine-readable OWL DL artifact, together with a populated, taxonomy-consistent decision-support layer and a SPARQL-based retrieval mechanism. It is deliberately not a claim to have built and deployed a complete recommender system. The artifact should be read as a semantic foundation: it fixes the vocabulary, the relational structure, and an initial suitability matrix $M : \mathcal{T} \times \mathcal{U} \times \mathcal{G} \rightarrow [0, 1]$ from which an executable recommender can be assembled. In this sense, the present release corresponds to the knowledge-formalization stage of a design-science process, while the construction of a production recommender—with a user interface, live data bindings, and empirically calibrated weights—is a subsequent stage. Framing the contribution this way clarifies why the evaluation in Section 5 is internal and artifact-centered, and why the heuristic, taxonomy-initialized weights are a starting point rather than a validated ranking model.

6.2. Validation Strategy: Internal Consistency and a Path to External Validation

The evaluation reported here establishes three properties: structural consistency (verified with the HermiT reasoner), competency-question satisfaction over the populated ABox, and internal consistency of the recommendation matrix against taxonomy-derived expectations. These are appropriate and necessary validation steps for an ontology artifact, but they are explicitly internal: they demonstrate that the artifact is logically sound and faithful to its design specification, not that its recommendations are optimal for real F&A practitioners. In particular, the numerical rankings are derived from the `taxonomyConsistencyScore` τ_{cons} and therefore inherit the ordering of the source taxonomy rather than being learned from evidence.

External validation would proceed along a concrete protocol. First, suitability weights would be elicited from a panel of practicing financial analysts, accountants, controllers, and CFOs through structured ranking tasks, and inter-rater agreement would be measured to replace τ_{cons} with expert-grounded values. Second, the recommendations produced by the ontology would be compared against expert judgments and, where feasible, against task-performance outcomes in controlled visualization-selection tasks. Third, the retrieval layer would be exercised on live F&A datasets so that the recommended techniques can be assessed for perceived usefulness and decision support. This staged protocol turns the internal-consistency evidence reported here into a testable claim of predictive validity, and directly addresses the current absence of independent user studies and of evidence-based justification for the numerical rankings.

6.3. Practical Applicability, Scalability, and Maintenance

Integrating the artifact into operational financial-analytics systems is feasible because the decision-support layer is exposed through standard SPARQL retrieval: an analytics application can issue a query with the current decision context (intention, user group, and interaction setting) and receive a ranked list of visualization techniques. The computational cost of retrieval is modest. The populated ABox contains 390 named individuals, including the full $3 \times 4 \times 13 = 156$ weight-matrix entries and the 156 corresponding ranked recommendations; queries over an ontology of this size execute in interactive time on a standard triple store, and the score-induced ranking is a simple ordering operation rather than a reasoning-intensive computation. Scalability along the context dimension is therefore governed by the size of the matrix, which grows multiplicatively with the number of intentions, user groups, and interaction settings rather than with the volume of the underlying financial data.

Maintenance is the principal ongoing cost. Because visualization techniques and interaction paradigms evolve, the class hierarchy and the weight matrix require periodic curation, and each added technique entails re-scoring across the affected decision contexts. We envisage this being controlled through lightweight governance—versioned releases of the OWL artifact, documented change lists, and, in future releases, SHACL constraints that check matrix population, cardinalities, and datatypes automatically so that incomplete or inconsistent updates are caught before publication. These measures keep the maintenance effort proportional to the number of changed techniques rather than to the size of the deployment.

6.4. Illustrative Operational Use Case

To make the intended operational role concrete, consider a quarterly financial-reporting dashboard for a mid-sized firm. A controller opens the dashboard with the intention of comparing segment performance across business units. The analytics front-end constructs the decision context—intention comparison, user group controller, and

an interaction setting supporting filtering and drill-down—and issues the corresponding SPARQL query against the ontology. The retrieval layer returns the taxonomy-consistent ranking of techniques for that context (for example, grouped bar and small-multiple layouts ranked above dense multivariate glyphs), and the dashboard instantiates the top-ranked technique while offering the next-ranked alternatives as one-click substitutions. Later, a CFO opens the same dashboard with the intention of obtaining a high-level overview; the context changes, a different query is issued, and the recommended encodings shift accordingly—mirroring the two contrasting worked scenarios in Section 4.9. This use case illustrates how the formalized knowledge is consumed at run time and where an executable recommender would attach, while also exposing the gaps that external validation must close: the ranking is only as good as the underlying weights, and the technique set must be kept current through the maintenance process described above.

7. Conclusions

This paper transformed the TaxUI&BV4FADA taxonomy into a machine-readable ontology framework for the selection of multidimensional and multivariate visualization techniques in financial and accounting (F&A) contexts. The resulting OWL DL artifact formalizes the four taxonomy-derived dimensions: visualization techniques, classification criteria including user intentions and visualization effects, interaction possibilities, and user groups (Sections 4.1–4.5).

On top of this taxonomy backbone, the artifact adds a decision-support layer composed of decision contexts, weight-matrix entries, and ranked recommendations, and a validation layer composed of competency questions, SPARQL validation queries, evaluation runs, and recommendation-evaluation records (Sections 4.7 and 4.8). At the implementation level, the artifact contains 198 classes, 18 object properties, 19 data properties, and 390 named individuals. These include the fully populated $3 \times 4 \times 13 = 156$ weight-matrix entries, the 156 corresponding ranked recommendation individuals covering all twelve canonical intention–user-group decision contexts, twelve internal recommendation-evaluation individuals gathered under `EvalRun_FullMatrix`, and eight competency questions paired with eight SPARQL validation-query individuals.

The evaluation reported here is internal and artifact-centered. Structural consistency was verified with the HermiT reasoner, which reported no inconsistencies and no unsatisfiable classes (Section 5.2). The eight encoded competency questions return the expected top-1, top- k , threshold-based, full-ranking, interaction-aware, matrix-cardinality, and per-context-completeness answers when their stored SPARQL queries are executed against the populated ABox (Section 5.3). The recommendation-matrix evaluation shows that the populated weight matrix is aligned with the taxonomy-derived expectations across the entire 3×4 context grid, with the taxonomy-consistency score τ_{cons} stored uniformly high on the corresponding `RecommendationEvaluation` individuals (Section 5.4). The two worked scenarios—a financial analyst testing a gross-margin hypothesis and a CFO seeking a quarterly data overview—demonstrate that the recommendation function is not globally constant: it discriminates between F&A roles and analytical intentions rather than producing a single global ordering of techniques (Section 4.9). Taken together, these results establish structural consistency, ABox queryability, and implementation fidelity between the design specification and the populated ontology artifact. They do not establish independent empirical recommendation accuracy, which would require external validation with F&A professionals (Section 5.5).

The artifact should therefore be read as a semantic foundation for recommendation retrieval rather than as a complete deployed recommender. Its main limitations are the heuristic, taxonomy-consistent initialization of the weight matrix; the absence of inde-

pendent expert elicitation and user-study validation; the use of the visualization-effect branch VTE as explanatory vocabulary rather than as an independent dimension of M ; the population of twelve canonical decision contexts rather than all $3 \times 4 \times 5 = 60$ possible interaction-aware combinations; the limited use of SWRL as anchor assertions rather than as a ranking algorithm; the absence of SHACL constraints for closed-world validation; and the fact that F&A semantics are currently documented through the external mapping σ rather than encoded through a separately imported financial/accounting ontology.

8. Future Work

Future work will proceed in four main directions. First, the taxonomy-initialized weight matrix should be replaced or calibrated with expert-elicited and empirically estimated weights, obtained through structured elicitation and user studies with practicing financial analysts, accountants, controllers, and CFOs. Second, the scoring model should be extended with additional dimensions—most notably the visualization-effect dimension and the data-characteristic extension M' discussed in Section 5.5. Third, future releases should add SHACL constraints and data-characteristic modeling for cardinality, completeness, datatype, and matrix-population validation. Finally, the ontology should be integrated with an executable visualization-recommendation prototype over live F&A datasets and linked, where appropriate, to financial-reporting and XBRL vocabularies, allowing the internal-consistency evidence reported here to be complemented by predictive validity and user-centered evaluation.

Author Contributions: Conceptualization, S.S. and S.L.; methodology, S.S. and S.L.; software, S.S. and S.L.; validation, S.L. and S.S.; formal analysis, S.S.; investigation, S.S.; resources, S.S. and S.L.; data curation, S.S. and S.L.; writing—original draft preparation, S.S.; writing—review and editing, S.L.; visualization, S.L.; supervision, S.S.; project administration, S.L. All authors have read and agreed to the published version of the manuscript.

Funding: This research received no external funding.

Institutional Review Board Statement: Not applicable.

Informed Consent Statement: Not applicable.

Data Availability Statement: The OWL ontology artifact (`FinAccOnt.ver-01.owl`), an equivalent Turtle serialization, and the eight SPARQL validation queries used in the internal validation are openly available in a public repository at <https://github.com/SuzanaLoshkovska/FinAccOnt> (accessed on 19 Aug 2026) to support reproducibility.

Acknowledgments: During the preparation of this manuscript we used large language models Claude Opus 4.6 and ChatGPT-5 Thinking for the purpose of generating figures, converting tables in Latex format, improving English language, and paper consistency check. The authors have reviewed and edited the output and take full responsibility for the content of this publication.

Conflicts of Interest: The authors declare no conflicts of interest.

Abbreviations

The following abbreviations are used in this manuscript:

ABox	assertional box
BFO	Basic Formal Ontology
CFO	chief financial officer
CQ	competency question
DCG	discounted cumulative gain
DL	description logic
EBITDA	earnings before interest, taxes, depreciation, and amortization

EP	evaluation process
F&A	financial and accounting
GAAP	Generally Accepted Accounting Principles
GDC	generically dependent continuant
IAO	Information Artifact Ontology
ICE	information content entity
IFRS	International Financial Reporting Standards
KPI	key performance indicator
MDMV	multidimensional and multivariate
MRR	mean reciprocal rank
nDCG	normalized discounted cumulative gain
OBI	Ontology for Biomedical Investigations
OWL	Web Ontology Language
P@k	precision at rank k
PP	planned process
RDF	Resource Description Framework
SDC	specifically dependent continuant
SHACL	Shapes Constraint Language
SPARQL	SPARQL Protocol and RDF Query Language
SWRL	Semantic Web Rule Language
TBox	terminological box
VQ	validation query
VTE	visualization technique effect
XBRL	eXtensible Business Reporting Language

Appendix A. Full Financial and Accounting Semantic Mapping

This appendix provides the element-level financial and accounting (F&A) interpretation of the taxonomy-derived ontology elements used in the recommendation layer. The mapping complements the compact dimension-level summary in Section 4.6, Table 4. The OWL identifiers are intentionally domain-neutral and are retained exactly as they appear in the submitted ontology artifact. Their F&A meaning is supplied through the external semantic mapping σ defined in Section 4.6. Consequently, the mappings in Table A1 support interpretation and documentation, but they do not introduce additional OWL axioms, imports, equivalence relations, or F&A-domain entailments.

Table A1. Element-level financial and accounting interpretation of taxonomy-derived OWL classes and ABox individuals referenced by the recommendation layer. Identifiers are retained as implemented in the OWL artifact; where legacy spelling is preserved, the intended reading is given in the interpretation column.

OWL Element	Role in Ontology	Financial and Accounting Interpretation
<i>User groups (UserGroups)—managerial and analytical roles in F&A organizations</i>		
TopManagers/ UG_TopManagers	User-group class/individual	C-suite, board-level, and CFO-level decision-makers concerned with strategic financial decision-making, regulatory reporting, risk exposure, and overall performance monitoring.
PlanControl	Role-specific subclass	Corporate planning and control, including budget planning, target setting, and high-level performance monitoring.
StratMan	Role-specific subclass	Strategic management of financial resources, investment priorities, capital allocation, and long-term financial positioning.

Table A1. Cont.

OWL Element	Role in Ontology	Financial and Accounting Interpretation
RegMan	Role-specific subclass	Regulatory and compliance reporting, including IFRS/GAAP-oriented reporting, disclosure monitoring, and supervisory reporting contexts.
TacticalManagers/ UG_TacticalManagers	User-group class/individual	Departmental, divisional, or business-unit financial managers who translate strategic goals into operational plans and monitor mid-level financial performance.
ManagByExep	Role-specific subclass	Management-by-exception over financial KPIs, especially variance reporting, threshold violations, and exception-driven control. The identifier is retained as implemented in the OWL file.
SectorPC	Role-specific subclass	Sectoral, departmental, or cost-center planning and control.
TacOperLevel	Role-specific subclass	Tactical–operational coordination between managerial planning and day-to-day financial execution.
OperativeManagers/ UG_OperativeManagers	User-group class/individual	Operational financial supervisors responsible for day-to-day process control, transaction monitoring, and anomaly handling.
StProcControl	Role-specific subclass	Statistical or process-oriented control over accounting transactions, operational costs, and routine financial flows.
ProbDetSolv	Role-specific subclass	Problem detection and resolution, including anomaly detection, exception handling, and corrective action on financial data.
OperManagers	Role-specific subclass	Routine operational management and supervision of financial or accounting processes.
AnalyticalStaff/ UG_AnalyticalStaff	User-group class/individual	Financial analysts, accountants, controllers, and internal auditors performing detailed analysis, hypothesis testing, and investigation of financial data.
ExceptionAn	Role-specific subclass	Exception analysis, including outlier transactions, unusual variance, fraud flags, and audit exceptions.
PerceptionAn	Role-specific subclass	Visual and perceptual analysis of financial charts, dashboards, and reports.
SpecAnTask	Role-specific subclass	Specialized or ad-hoc analytical tasks, including targeted investigations and non-routine financial analysis.
<i>User intentions (UserIntention)—F&A analytical goals</i>		
UL_DO/Crit_UL_DO	Intention class/ criterion individual	Data overview: periodic reporting, KPI monitoring, dashboard inspection, and summarization of financial performance.
UL_HC/Crit_UL_HC	Intention class/ criterion individual	Hypothesis confirmation: variance analysis, ratio testing, gross-margin hypothesis checking, and validation of expected financial relationships.
UL_DID_MND/ Crit_UL_DID_MND	Intention class/ criterion individual	Delve into data and make new decisions: exploratory investigation, capital-allocation analysis, strategic-pivot analysis, what-if reasoning, and decision discovery.
<i>Visualization effects (VTE)—explanatory perceptual and interpretive goals</i>		
Visibility	VTE first-level class	Clear exposure of accounting patterns, trends, deviations, outliers, and visually salient financial signals.
VisAnAg	VTE leaf class	Visibility of analytical aggregates or aggregated financial summaries.
VisObjSel	VTE leaf class	Visibility of selected financial objects, such as accounts, KPIs, transactions, periods, or business units.

Table A1. Cont.

OWL Element	Role in Ontology	Financial and Accounting Interpretation
VisRelation	VTE leaf class	Visibility of relationships among financial variables, indicators, accounts, or organizational units.
VisScr	VTE leaf class	Visibility of screen-level or report-level financial information needed for overview and inspection.
Interpretability	VTE first-level class	Cognitive ease of reading, explaining, and interpreting financial reports and visual displays.
IntAgg	VTE leaf class	Interpretability of aggregated financial measures and summaries.
IntLevData	VTE leaf class	Interpretability across levels of financial data, such as account, cost-center, department, or organizational hierarchy.
IntRelation	VTE leaf class	Interpretability of relationships among financial indicators, ratios, accounts, or reporting categories.
Insight_in_data	VTE first-level class	Discovery of new financial dependencies, structures, clusters, correlations, or decision-relevant patterns.
InsDataAss	VTE leaf class	Insight into associations among financial variables or indicators.
InsDataClass	VTE leaf class	Insight into classification or grouping of financial entities, accounts, customers, branches, or transactions.
InsDataClu	VTE leaf class	Insight into clusters within financial or accounting datasets.
InsDataCor	VTE leaf class	Insight into correlations among financial ratios, KPIs, or transaction attributes.
InsDataPoss	VTE leaf class	Insight into possible data-driven financial explanations, scenarios, or decision alternatives.
InsDataPr	VTE leaf class	Insight into financial patterns, profiles, or predictive tendencies.
<i>Interaction possibilities (InteractionPoss)—analyst actions on F&A data</i>		
SelectSubset/ Int_SelectSubset	Interaction class/individual	Period, segment, region, account group, product group, or business-unit selection.
SelectDataAttributes/ Int_SelectDataAttributes	Interaction class/individual	Selection of accounts, KPIs, ratios, dimensions, or attributes such as revenue, margin, cost, liquidity, working capital, or transaction type.
IPFilter/Int_IPFilter	Interaction class/individual	Filtering by amount, threshold, accounting category, period, segment, variance level, or exception condition.
SFIL1	Interaction subclass	Single-filter operation over an F&A dataset.
SFILn	Interaction subclass	Multiple-filter operation over several financial attributes, categories, or periods.
SOBJn	Interaction subclass	Selection of multiple financial objects, such as accounts, transactions, KPIs, or reporting units.
STFIL	Interaction subclass	Stepwise or temporal filtering, for example filtering by reporting period, quarter, month, or time window.
IPInt/Int_IPInt	Interaction class/individual	Interactive exploration, drill-down, navigation, and manipulation of financial or accounting views.
SDIS	Interaction subclass	Display or visual-detail interaction in a financial visualization.
SIF	Interaction subclass	Interactive filtering or focus operation.
SLB	Interaction subclass	Label-based interaction, such as showing account, KPI, branch, or transaction labels.
SNO	Interaction subclass	Node or object selection, especially in network, hierarchy, or object-based financial views.
SZOOM	Interaction subclass	Zooming into a financial view, time interval, cluster, branch, or region of interest.

Table A1. Cont.

OWL Element	Role in Ontology	Financial and Accounting Interpretation
PossibilityInterWithData/Int_PossibilityInterWithData	Interaction class/individual	General interaction with financial data, including sorting, brushing, zooming, navigating, and changing visual focus.
<i>Visualization techniques (MDMV_techniques)—F&A visual representations</i>		
Dashboards/Tech_Dashboards	Technique class/individual	Financial KPI dashboards for overview of revenue, gross margin, EBITDA, liquidity, profitability, variance, budget status, and other managerial indicators.
histograms/Tech_histograms	Dashboard sub-technique	Distribution views for transaction amounts, expense categories, account balances, error frequencies, or audit exceptions.
scatter_plots/Tech_scatter_plots	Dashboard sub-technique	Pairwise financial relationship analysis, such as revenue vs. margin, cost vs. volume, or risk vs. return.
line_plots/Tech_line_plots	Dashboard sub-technique	Trend displays for time-dependent financial measures such as sales, costs, cash flow, budget variance, or profitability.
bubble_plots/Tech_bubble_plots	Dashboard sub-technique	Bubble plots for multivariate financial comparisons, where marker size may encode volume, exposure, revenue, or risk.
time_series/Tech_time_series	Dashboard sub-technique	Temporal financial reporting and monitoring across periods, quarters, months, or fiscal years.
Scatterplots/Tech_Scatterplots	Technique class/individual	Cross-account, cross-period, or cross-entity scatter analysis for financial relationships and anomalies.
HierarScatterplots/Tech_HierarScatterplots	Hierarchical variant	Scatterplot-based analysis across hierarchical financial structures such as regions, branches, product groups, or account categories.
ParCor/Tech_ParCor	Technique class/individual	Parallel-coordinates analysis of multidimensional financial ratios, KPIs, performance indicators, and transaction attributes.
HierarParCor/Tech_HierarParCor	Hierarchical variant	Parallel-coordinates analysis for nested F&A structures, such as account hierarchies, departmental hierarchies, or portfolio groupings.
Pixel_oriented_display/Tech_Pixel_oriented_display	Technique class/individual	Pixel-oriented visualization for very large transaction tables, ledger entries, audit trails, or high-volume accounting records.
HierarPixOriented_display/Tech_HierarPixOriented_display	Hierarchical variant	Pixel-oriented display over grouped or nested financial structures, such as account groups, cost centers, or organizational units.
Glyphs/Tech_Glyphs	Technique class/individual	Glyph-based summaries where one glyph may represent an account, customer, branch, project, portfolio item, or business unit.
HierarGlyphs/Tech_HierarGlyphs	Hierarchical variant	Hierarchical glyph-based summaries for nested portfolios, branches, accounts, or organizational structures.
DimensionalStacking/Tech_DimensionalStacking	Technique class/individual	Dimensional stacking for multilevel financial dimensions, including chart-of-accounts hierarchies, product structures, regional structures, or budget categories.
HierarDimenStacking/Tech_HierarDimenStacking	Hierarchical variant	Hierarchical dimensional stacking for deeply nested financial and accounting structures.
Nodes_and_Links/Tech_Nodes_and_Links	Technique class/individual	Network views for intercompany flows, ownership structures, transaction networks, fraud-link analysis, supplier/customer relationships, or audit trails.
Polar_coordinates/Tech_Polar_coordinates	Technique class/individual	Cyclic or seasonal financial pattern displays, such as monthly seasonality, fiscal cycles, periodic costs, or recurring revenue patterns.

Appendix B. Validation-Query SPARQL Listings

The eight SPARQL validation queries VQ_CQ1–VQ_CQ8 are stored in the OWL file as queryText data-property values on the corresponding ValidationQuery individuals and executed against the populated ABox during competency-question validation (Section 5.3). All eight queries return the result stored as expectedAnswerText on the corresponding CompetencyQuestion individual, so all eight validationPassed flags evaluate to true (Table 11). The listings below group them by facet.

Single-best-answer retrieval (VQ_CQ1, VQ_CQ2).

Both queries locate a CTX_* decision-context individual and read its recommendsTechnique link directly.

```
PREFIX ex: <http://webprotege.stanford.edu/project/MMwphrtbm7u64GTgmDamg#>
```

```
# VQ_CQ1 --> Tech_Dashboards
SELECT ?technique WHERE {
  ex:CTX_UL_DO_TopManagers ex:recommendsTechnique ?technique .
}
```

```
# VQ_CQ2 --> Tech_ParCor
SELECT ?technique WHERE {
  ex:CTX_UL_HC_AnalyticalStaff ex:recommendsTechnique ?technique .
}
```

Top-*k* retrieval (VQ_CQ3).

Walks the Recommendation individuals attached to a context and filters by rank.

```
PREFIX ex: <http://webprotege.stanford.edu/project/MMwphrtbm7u64GTgmDamg#>
```

```
# VQ_CQ3 --> Tech_DimensionalStacking, Tech_HierarDimenStacking,
#           Tech_Nodes_and_Links
SELECT ?technique ?rank WHERE {
  ?r a ex:Recommendation ;
    ex:hasContext ex:CTX_UL_DID_MND_TacticalManagers ;
    ex:hasVisualizationTechnique ?technique ;
    ex:hasRank ?rank .
  FILTER(?rank <= 3)
}
ORDER BY ?rank
```

Interaction-aware retrieval (VQ_CQ4).

Reads the considersInteraction link of a decision context.

```
PREFIX ex: <http://webprotege.stanford.edu/project/MMwphrtbm7u64GTgmDamg#>
```

```
# VQ_CQ4 --> Int_SelectDataAttributes
SELECT ?interaction WHERE {
  ex:CTX_UL_HC_AnalyticalStaff ex:considersInteraction ?interaction .
}
```

Matrix-level cardinality (VQ_CQ5).

A counting aggregate over WeightMatrixEntry individuals.

PREFIX ex: <http://webprotege.stanford.edu/project/MMwphrtbm7u64GTgmDamg#>

```
# VQ_CQ5 --> 156
SELECT (COUNT(?w) AS ?count) WHERE {
  ?w a ex:WeightMatrixEntry .
}
```

Threshold-based retrieval (VQ_CQ6).

A score-filtered selection over Recommendation individuals attached to a context.

PREFIX ex: <http://webprotege.stanford.edu/project/MMwphrtbm7u64GTgmDamg#>

```
# VQ_CQ6 --> Tech_ParCor, Tech_HierarParCor, Tech_Scatterplots,
#           Tech_HierarScatterplots, Tech_Pixel_oriented_display
SELECT ?technique ?score WHERE {
  ?r a ex:Recommendation ;
    ex:hasContext ex:CTX_UL_HC_AnalyticalStaff ;
    ex:hasVisualizationTechnique ?technique ;
    ex:recommendationScore ?score .
  FILTER(?score >= 0.85)
}
ORDER BY DESC(?score)
```

Full-ranking retrieval (VQ_CQ7).

Returns the complete ranked list with both score and rank for a single context.

PREFIX ex: <http://webprotege.stanford.edu/project/MMwphrtbm7u64GTgmDamg#>

```
# VQ_CQ7 --> 13 ranked rows for CTX_UL_DO_TopManagers
SELECT ?technique ?score ?rank WHERE {
  ?r a ex:Recommendation ;
    ex:hasContext ex:CTX_UL_DO_TopManagers ;
    ex:hasVisualizationTechnique ?technique ;
    ex:recommendationScore ?score ;
    ex:hasRank ?rank .
}
ORDER BY ?rank
```

Per-context completeness (VQ_CQ8).

A diagnostic GROUP BY/HAVING that returns only contexts whose recommendation count differs from 13. An empty result set confirms that all twelve full-matrix CTX_* contexts hold exactly thirteen recommendations each.

PREFIX ex: <http://webprotege.stanford.edu/project/MMwphrtbm7u64GTgmDamg#>

```
# VQ_CQ8 --> empty result set
SELECT ?ctx (COUNT(?r) AS ?cnt) WHERE {
  ?r a ex:Recommendation ;
    ex:hasContext ?ctx .
  ?ctx a ex:DecisionContext .
  FILTER(STRSTARTS(STR(?ctx), STR(ex:CTX_)))
}
```

```
GROUP BY ?ctx
HAVING (COUNT(?r) != 13)
```

These eight listings together cover the six facets of the competency-question layer (single-best-answer, top-*k*, interaction-aware, matrix-level cardinality, threshold-based, full-ranking, and per-context completeness retrieval; see Section 5.3). Each query in the OWL file is identified by its queryText value on the corresponding VQ_CQk individual and can be re-executed against the populated graph without any modification to the artifact.

References

1. Gruber, T. Ontology. In *Encyclopedia of Database Systems*; Liu, L., Özsu, M.T., Eds.; Springer: New York, NY, USA, 2018; pp. 2574–2576. [\[CrossRef\]](#)
2. Borst, W.N. Construction of Engineering Ontologies for Knowledge Sharing and Reuse. Ph.D. Thesis, University of Twente, Enschede, The Netherlands, 1997. [\[CrossRef\]](#)
3. Noy, N.F.; McGuinness, D.L. *Ontology Development 101: A Guide to Creating Your First Ontology*; Technical Report KSL-01-05; Stanford Knowledge Systems Laboratory: Stanford, CA, USA, 2001.
4. Uschold, M.; Gruninger, M. Ontologies: Principles, Methods and Applications. *Knowl. Eng. Rev.* **1996**, *11*, 93–136. [\[CrossRef\]](#)
5. Inselberg, A.; Dimsdale, B. Parallel Coordinates: A Tool for Visualizing Multi-Dimensional Geometry. In Proceedings of the First IEEE Conference on Visualization (Visualization '90), San Francisco, CA, USA, 23–26 October 1990; pp. 361–378. [\[CrossRef\]](#)
6. Carr, D.B.; Littlefield, R.J.; Nicholson, W.L.; Littlefield, J.S. Scatterplot Matrix Techniques for Large N. *J. Am. Stat. Assoc.* **1987**, *82*, 424–436. [\[CrossRef\]](#)
7. Chernoff, H. The Use of Faces to Represent Points in *k*-Dimensional Space Graphically. *J. Am. Stat. Assoc.* **1973**, *68*, 361–368. [\[CrossRef\]](#)
8. Ward, M.O. A Taxonomy of Glyph Placement Strategies for Multidimensional Data Visualization. *Inf. Vis.* **2002**, *1*, 194–210. [\[CrossRef\]](#)
9. LeBlanc, J.; Ward, M.O.; Wittels, N. Exploring N-Dimensional Databases. In Proceedings of the First IEEE Conference on Visualization (Visualization '90), San Francisco, CA, USA, 23–26 October 1990; pp. 230–237. [\[CrossRef\]](#)
10. Keim, D.A. Pixel-Oriented Visualization Techniques for Exploring Very Large Data Bases. *J. Comput. Graph. Stat.* **1996**, *5*, 58–77. [\[CrossRef\]](#)
11. Pickett, R.M.; Grinstein, G.G. Iconographic Displays for Visualizing Multidimensional Data. In Proceedings of the 1988 IEEE International Conference on Systems, Man, and Cybernetics, Beijing, China, 8–12 August 1988; Volume 1, pp. 514–519. [\[CrossRef\]](#)
12. Shneiderman, B. Tree Visualization with Tree-Maps: 2-D Space-Filling Approach. *ACM Trans. Graph.* **1992**, *11*, 92–99. [\[CrossRef\]](#)
13. Kruskal, J.B.; Landwehr, J.M. Icicle Plots: Better Displays for Hierarchical Clustering. *Am. Stat.* **1983**, *37*, 162–168. [\[CrossRef\]](#)
14. Shneiderman, B. The Eyes Have It: A Task by Data Type Taxonomy for Information Visualizations. In Proceedings of the IEEE Symposium on Visual Languages, Boulder, CO, USA, 3–6 September 1996; pp. 336–343. [\[CrossRef\]](#)
15. Card, S.K.; Mackinlay, J.D.; Shneiderman, B. *Readings in Information Visualization: Using Vision to Think*; Morgan Kaufmann: San Francisco, CA, USA, 1999.
16. Larrea, M.L.; Martig, S.; Castro, S.M. A Semantics-based visualization building process. In Proceedings of the XVI Congreso Argentino de Ciencias de la Computación (CACIC 2010), Morón, Argentina, 18–22 October 2010; pp. 465–474.
17. Shu, G.; Avis, N.J.; Rana, O.F. Bringing semantics to visualization services. *Adv. Eng. Softw.* **2008**, *39*, 514–520. [\[CrossRef\]](#)
18. Savoska, S.; Loshkovska, S. Taxonomy of User Intention and Benefits of Visualization for Financial and Accounting Data Analysis. In *ICT Innovations 2013*; Trajkovic, V., Anastas, M., Eds.; Springer: Berlin/Heidelberg, Germany, 2014; pp. 197–207.
19. Savoska, S.; Loskovska, S. Evaluation of Taxonomy of User Intention and Benefits of Visualization for Financial and Accounting Data Analysis. In Proceedings of the 7th International Conference on Information Systems and Grid Technologies (ISGT 2013), Sofia, Bulgaria, 31 May–1 June 2013.
20. Bechhofer, S.; van Harmelen, F.; Hendler, J.; Horrocks, I.; McGuinness, D.L.; Patel-Schneider, P.F.; Stein, L.A. *OWL Web Ontology Language Reference*; Technical Report; W3C Recommendation; World Wide Web Consortium (W3C): Cambridge, MA, USA, 2004.
21. Musen, M.A. The Protégé Project: A Look Back and a Look Forward. *AI Matters* **2015**, *1*, 4–12. [\[CrossRef\]](#) [\[PubMed\]](#)
22. Bertin, J. *Semiology of Graphics: Diagrams, Networks, Maps*; University of Wisconsin Press: Madison, WI, USA, 1983.
23. McCormick, B.H. Visualization in scientific computing. *SIGBIO Newsl.* **1988**, *10*, 15–21. [\[CrossRef\]](#)
24. Tufte, E.R. *The Visual Display of Quantitative Information*, 2nd ed.; Graphics Press: Cheshire, CT, USA, 2001.
25. Mackinlay, J.D. Automating the Design of Graphical Presentations of Relational Information. *ACM Trans. Graph.* **1986**, *5*, 110–141. [\[CrossRef\]](#)
26. Spence, R. *Information Visualization: Design for Interaction*, 2nd ed.; Prentice Hall: Harlow, UK, 2007.

27. Tory, M.; Möller, T. Rethinking Visualization: A High-Level Taxonomy. In Proceedings of the IEEE Symposium on Information Visualization, INFO VIS, Austin, TX, USA, 10–12 October 2004; pp. 151–158. [\[CrossRef\]](#)
28. Katifori, A.; Halatsis, C.; Lepouras, G.; Vassilakis, C.; Giannopoulou, E. Ontology Visualization Methods—A Survey. *ACM Comput. Surv.* **2007**, *39*, 10. [\[CrossRef\]](#)
29. Brodlie, K.W.; Duke, D.J.; Duce, D.A. Building an Ontology of Visualization. In Proceedings of the IEEE Visualization Conference (VIS 2004)—Poster, Los Alamitos, CA, USA, 10–15 October 2004; pp. 1–7. [\[CrossRef\]](#)
30. Polowinski, J.; Voigt, M. VISO: A Shared, Formal Knowledge Base as a Foundation for Semi-Automatic InfoVis Systems. In Proceedings of the CHI '13 Extended Abstracts on Human Factors in Computing Systems, Paris, France, 27 April–2 May 2013; pp. 1791–1796. [\[CrossRef\]](#)
31. Del Rio, N.; da Silva, P.P. *VisKo: Semantic Web Support for Information and Scientific Visualization*; Technical Report UTEP-CS-11-58; University of Texas at El Paso: El Paso, TX, USA, 2011.
32. Sacha, D.; Kraus, M.; Keim, D.A.; Chen, M. VIS4ML: An Ontology for Visual Analytics Assisted Machine Learning. *IEEE Trans. Vis. Comput. Graph.* **2019**, *25*, 385–395. [\[CrossRef\]](#) [\[PubMed\]](#)
33. Fabbri, R.; de Oliveira, M.C.F. Audiovisual Analytics Vocabulary and Ontology (AAVO): Initial Core and Example Expansion. *arXiv* **2017**, arXiv:1710.09954. [\[CrossRef\]](#)
34. Loshkovska, S.; Panov, P. Foundations for a Generic Ontology for Visualization: A Comprehensive Survey. *Information* **2025**, *16*, 915. [\[CrossRef\]](#)
35. Few, S. *Information Dashboard Design: Displaying Data for At-a-Glance Monitoring*, 2nd ed.; Analytics Press: Burlingame, CA, USA, 2013.
36. Eckerson, W.W. *Performance Dashboards: Measuring, Monitoring, and Managing Your Business*, 2nd ed.; Wiley: Hoboken, NJ, USA, 2010.
37. Negash, S.; Gray, P. Business Intelligence. In *Handbook on Decision Support Systems 2: Variations*; Springer: Berlin/Heidelberg, Germany, 2008; pp. 175–193. [\[CrossRef\]](#)
38. Korczak, J.; Dudycz, H.; Dyczkowski, M. Intelligent Dashboard for SME Managers. Architecture and Functions. In Proceedings of the Federated Conference on Computer Science and Information Systems (FedCSIS 2012), Wroclaw, Poland, 9–12 September 2012; pp. 1003–1007.
39. Korczak, J.; Dudycz, H.; Dyczkowski, M. Design of Financial Knowledge in Dashboard for SME Managers. In Proceedings of the Federated Conference on Computer Science and Information Systems (FedCSIS 2013), Krakow, Poland, 8–11 September 2013; pp. 1111–1118.
40. Dyczkowski, M.; Korczak, J.; Dudycz, H. Multi-criteria evaluation of the intelligent dashboard for SME managers based on scorecard framework. In Proceedings of the 2014 Federated Conference on Computer Science and Information Systems, Warsaw, Poland, 7–10 September 2014; pp. 1147–1155. [\[CrossRef\]](#)
41. Dudycz, H.; Korczak, J. Conceptual design of financial ontology. In Proceedings of the 2015 Federated Conference on Computer Science and Information Systems (FedCSIS), Łódź, Poland, 13–16 September 2015; pp. 1505–1511. [\[CrossRef\]](#)
42. Dudycz, H. Application of Semantic Network Visualization as a Managerial Support Instrument in Financial Analyses. *Online J. Appl. Knowl. Manag.* **2017**, *5*, 112–128. [\[CrossRef\]](#)
43. Korczak, J.; Dudycz, H.; Nita, B.; Oleksyk, P. Towards Process-Oriented Ontology for Financial Analysis. In Proceedings of the Federated Conference on Computer Science and Information Systems (FedCSIS 2017), Prague, Czech Republic, 3–6 September 2017; pp. 981–987. [\[CrossRef\]](#)
44. Jedrzejek, C.; Falkowski, M.; Smolenski, M. Link Analysis of Fuel Laundering Scams and Implications of Results for Scheme Understanding and Prosecutor Strategy. *Front. Artif. Intell. Appl.* **2009**, *205*, 79–88. [\[CrossRef\]](#)
45. Abrouk, L.; Chergui, H.; Ahaggach, H. OntoFiC: An ontology for financial fraud detection and customer behavior modeling. In Proceedings of the ASONAM '23: 2023 IEEE/ACM International Conference on Advances in Social Networks Analysis and Mining, New York, NY, USA, 6–9 November 2024; pp. 509–513. [\[CrossRef\]](#)
46. Marshall, B.; Mortenson, K.; Bourne, A.; Price, K. Visualizing Basic Accounting Flows: Does XBRL + Model + Animation = Understanding? *Int. J. Digit. Account. Res.* **2010**, *10*, 27–54. [\[CrossRef\]](#) [\[PubMed\]](#)
47. Bonson, E.; Escobar, T.; Flores, F. BaselMapper: An experimental tool for managing operational risk in banks. *Int. J. Metadata Semant. Ontol.* **2011**, *6*, 1–9. [\[CrossRef\]](#)
48. Hu, B.; Naseer, A.; Rodrigues, E.M.; Vandenbussche, P.Y.; Goto, M. Linked Data for Financial Reporting. In Proceedings of the Fourth International Workshop on Consuming Linked Data, Collocated with the 12th International Semantic Web Conference, Sydney, Australia, 22 October 2013.
49. O'Riain, S.; Curry, E.; Harth, A. XBRL and open data for global financial ecosystems: A linked data approach. *Int. J. Account. Inf. Syst.* **2012**, *13*, 141–162. [\[CrossRef\]](#)
50. Hevner, A.R.; March, S.T.; Park, J.; Ram, S. Design science in information systems research. *MIS Q.* **2004**, *28*, 75–105. [\[CrossRef\]](#)
51. Peffers, K.; Tuunanen, T.; Rothenberger, M.A.; Chatterjee, S. A Design Science Research Methodology for Information Systems Research. *J. Manag. Inf. Syst.* **2008**, *24*, 45–77. [\[CrossRef\]](#)

52. Cyganiak, R.; Wood, D.; Lanthaler, M. *RDF 1.1 Concepts and Abstract Syntax*; W3c recommendation; World Wide Web Consortium (W3C): Cambridge, MA, USA, 2014.
53. Glimm, B.; Horrocks, I.; Motik, B.; Stoilos, G.; Wang, Z. HermiT: An OWL 2 Reasoner. *J. Autom. Reason.* **2014**, *53*, 245–269. [[CrossRef](#)]
54. Sirin, E.; Parsia, B.; Grau, B.C.; Kalyanpur, A.; Katz, Y. Pellet: A Practical OWL-DL Reasoner. *J. Web Semant.* **2007**, *5*, 51–53. [[CrossRef](#)]
55. Kazakov, Y.; Krötzsch, M.; Simancik, F. The Incredible ELK: From Polynomial Procedures to Efficient Reasoning with EL Ontologies. *J. Autom. Reason.* **2014**, *53*, 1–61. [[CrossRef](#)]

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.