

A Data-Driven Decision Support Framework for Business Intelligence Systems Leveraging BERT-Enhanced LLMs

Alaa Abid Muslam Abid Ali*

Cyber Security Department, Faculty of Computer Science and Information Technology Al-Qadisiyah University, Al-Qadisiyah, Iraq

Abstract Many techniques were introduced to intimate the gap existing between raw organizational data of business intelligence systems and practical decision support. One of these techniques is to integrate the large language models (LLMs) into enterprise business intelligence (BI) architectures. The available BI frameworks suffer from bi-direction drawbacks. These drawbacks represented by treat natural language query interfaces as peripheral additions or by failing in providing empirically confirmed and decision-quality outcomes. This will put practitioners without any principled basis for architectural investment. This paper addresses both limitations. We present a three-layer BI-Enabled Decision Quality (BIDQ) framework that formalizes the interdependencies among data readiness, analytical capability, and decision integration, and we introduce a BERT-based query understanding module trained on 180,000 enterprise query–response pairs drawn from five production systems (SAP ERP, Salesforce CRM, Oracle SCM, customer service ticketing, and financial reporting). The BERT model achieves 86.5% test accuracy, an F1-score of 86%, and a mean inference latency of 45 ms a 5.5-percentage-point accuracy improvement over the best LSTM baseline ($p < 0.001$, Cohen’s $d = 0.42$). Training converged in 72 GPU-hours on eight NVIDIA A100 GPUs using AdamW optimization with linear warmup and decay. To ground the framework in measurable outcomes, we conducted a twelve-month field study across five industry sectors (retail, banking, healthcare, manufacturing, and logistics), finding an average 26-percentage-point improvement in decision quality and an 89% reduction in decision cycle time following BI deployment. Financial modelling indicates a typical break-even point at approximately month 12 of deployment. These results determine the quantitative criterions that can guide the organizations in formalization and structuring the investments and measuring of LLM-augmented BI.

Keywords Large Language Models, BERT, Business Intelligence, Query Understanding, Decision Support Systems, Model Training, Transformer Architecture, Analytical Maturity, Enterprise Analytics

DOI: 10.19139/soic-2310-5070-3287

1. Introduction

One of the biggest problems with business nowadays is that they are inundated and intolerable with the amount of data they have, but do not use it effectively for their decisions. Think about it this way; companies, are flooded with massive amounts of data from various sources such as transactions, sensors, customer interactions and market trends. Yet, as they approach the decision-making process, they realize how often they do not have access to the right piece of information at the right moment. Instead, they are either relying on an antiquated report, a delayed update or their memory. Not only is this approach against the common sense but also completely inefficient, especially since Business Intelligence systems were around the corner with a solution to solve this issue in mind. They’re designed to collect all the different types of data, process them and deliver insight in an easily understandable and actionable way for those making decisions. I think the real issue is that companies are not getting the maximum out of these systems. This state is indigenous to practitioners of the field who ignore all the data they have in aiming to make and lead rational decision. Consequently, the corporations neglect opportunities

*Correspondence to: Alaa Abid Muslam Abid Ali (Email: alaa.abidmuslam@qu.edu.iq). Cyber Security Department, Faculty of Computer Science and Information Technology, Al-Qadisiyah University, Al-Qadisiyah, Iraq.

to build better business decisions based on data, and train themselves for enhanced results that lead them to more successful at the end.

A new dimension was introduced to this effort by Transformer-based large language models. While earlier-generation BI systems required users to write queries in structured query languages or drill through pre-defined report hierarchies, LLM-augmented interfaces enable decision-makers to read data repositories as they speak: asking questions in natural language that mirrors how they think about problems and receiving answers synthesized from the entire range of available organizational data [3, 4]. BERT (Bidirectional Encoder Representations from Transformers) [5] in particular has performed well at enterprise query understanding, as it is designed to capture bidirectional contextual relationships within a query — something simpler approaches struggle with due to the potential for semantic ambiguity.

Even with substantial work on BI system design and on adapting transformer models to specific domains, two gaps in the literature remain. For starters, although existing BI frameworks describe system components from an architectural perspective, they only indirectly convey the manner in which a technical architectural choice propagates into its decision quality outcomes [6, 7]. Second, the performance of LLM-augmented BI systems usually done in the labs rather than actual deployments of organizations using benchmark datasets [8]. This consequently throw doubts with the ecological validity of such results. The present work bridges both gaps.

The paper makes the following research contributions: A three-tier BI-Enabled Decision Quality (BIDQ) framework that leverages the sequential structure of dependence between data readiness, analytical capability, and decision integration and is intended to be both definitive (in terms of how these factors fit together) as well as diagnostic for guiding users in identifying the binding constraint within any specific organizational context.

A natural language query understanding module based on BERT, trained on 180,000 enterprise-level query-response pairs across five decision domains with complete documentation on training infrastructure, optimization strategy, convergence behavior and metrics during deployment in production. A long study across year different like retail, banking, healthcare, industries manufacturing, and logistics provides a basis how well decisions for measuring are made, it takes to make them, and the how long financial benefits of using. This study gives Business Intelligence tools us standards against, so we to compare future projects can see how Novelty and positioning relative to prior work. The contribution of this paper is not the use of BERT or weak supervision in isolation, both of which are established, but their combination into a deployment-grounded decision-quality argument that prior work does not make. Existing BI-maturity and analytical-maturity frameworks, including the Gartner and Sharda staircases and the descriptive-to-prescriptive progression of Chen et al., describe what an organization can do at each level but neither specify the binding prerequisites that gate movement between levels nor connect a specific architectural choice to a measured decision-quality outcome; BIDQ adds exactly this by formalizing the multiplicative dependence among data readiness, analytical capability, and decision integration and by making the binding constraint diagnosable. Relative to prior text-to-SQL and BERT-based query-understanding work, such as Devlin et al. and Min et al., which is evaluated almost entirely on public benchmarks under laboratory conditions, our novelty is empirical rather than architectural: we report training, weak-supervision labeling, and CPU-served inference on 180,000 enterprise query-response pairs from five production systems, together with a twelve-month field study, producing an ecologically valid measurement of an existing technique in a setting where such measurements are scarce. We accordingly frame the technical module as an application whose value is the in-situ evidence it produces, and we have narrowed the contribution claims to this scope.

The rest of this paper is laid out in the following way. First, we'll take a look at what other people have done in the areas of business intelligence architecture, analytical maturity modeling, and using large language models to understand queries. Then, we'll introduce the BIDQ framework. After that, we'll describe how the whole business intelligence and big data system works from start to finish. Upcoming, we will formulate a model to assess the maturity of enterprise analytics. Also, a discussion with full details for the part of the system related with BERT will done to understand queries. This followed by sharing the results of our established experiments addition to field study we conducted. Next, we will conduct a discussion that review how this model can be used in different sectors and the challenges we encountered and also planned future research directions. Finally, we will introduce a wide summary of all experiments and results.

2. Related Work

2.1. Business Intelligence Architecture and Decision Support

The conceptual foundations of decision support systems can be traced back to Gore and Scott Morton [9]. They distinguished between structured, semi-structured, and unstructured decisions. The researchers based their argument on the fact that information requirements differ between these categories. Modern business intelligence architectures embody this vision through a pipeline that transfers data from source systems through integration, storage, and analysis layers to the decision support interface [1, 10]. Chen et al. [2] provided an influential synthesis of BI and analytics research, identifying a progression from descriptive statistics toward predictive and prescriptive analytics as the defining trajectory of the field. However, this body of work predominantly treats BI components in isolation and offers limited formal analysis of how component-level constraints interact to limit system-level decision outcomes a gap addressed by the BIDQ framework proposed here.

To get the most out of business intelligence, you need good quality data - it's a major obstacle to getting real value from it. When designing a data warehouse, it's crucial to tackle the issue of different data sources in a systematic way, rather than trying to clean up the data as an afterthought, as Inmon suggests. Recent studies, like the one by Vidgen et al., have shown that a company's ability to adopt business intelligence has more to do with its culture than its technical abilities.

For example, if senior managers make decisions based on data, the extent of their use of data is a factor. This implies that having the correct architecture is not sufficient to enhance the quality of decision making. The key is to build a culture of consistent use of data, particularly by the leaders. If the senior managers use data to make decisions, then the tone will set for the rest of the organisation. The cultural change may make a huge difference to the success of business intelligence projects. From the outset, data quality becomes a top priority, and once it's addressed, it becomes easier for organizations to create a culture that values data-driven decision-making and reap the benefits of their business intelligence initiatives. Having the right tools and technology is not enough, it's also about fostering a culture that embraces data and leverages it to make decisions. Actually, leaders that focus on data-driven decision making can spread the influence throughout the organisation, resulting in improved decision making and more effective business intelligence use. To do so, it is crucial to foster a culture that cherishes data and utilizes it in all aspects, not merely technology to fix the issue. This will help organizations maximize their use of business intelligence and make informed business decisions.

2.2. Analytical Maturity Frameworks

Companies' journey in how they are building the capacity to use data and analytics is visible. This track includes four phases as illustrated in some models, such as the one from Gartner: companies can describe what is happening, they can diagnose problems, they can predict what will happen, and then they can prescribe what to do. Another model, by Sharda and colleagues, divides it into five phases, each applicable to more advanced technology and skills. Companies that are further down this continuum enjoy improved financial performance compared to their counterparts, said Davenport and Harris. These models, however, don't give us any insight into what companies must have in place to achieve a step up: the right data infrastructure, the right employees, the right culture. The model we present later in this paper tries to fill this gap by making clear what's needed at each level to make progress.

2.3. LLMs and Natural Language Interfaces for Enterprise Analytics

Recent, the applied of transformer-based language models to undertaking query has enticing. Devlin et al. [5] introduced BERT and demonstrated state-of-the-art performance across eleven natural language understanding tasks through bidirectional pre-training on large text corpora. Subsequent work fine-tuned BERT and its derivatives for text-to-SQL generation [16], a task closely related to the query understanding problem addressed in this paper. Min et al. [17] demonstrated that fine-tuned BERT models outperform LSTM-based approaches on question answering benchmarks by capturing longer-range contextual dependencies that recurrent architectures miss.

Most research on using natural language to ask databases questions has been done in controlled environments. But

what happens when you take this technology into the real world, where the language is specific to the industry and the system has to make quick decisions? This is a big gap in our knowledge, and it's what this study aims to fill. The researchers looked at how well a system can be trained and used in five different business systems, and how it affects the decisions made in real-world situations. Researchers focused on situations where the system needs to understand specialized terminology used in a specific field, such as a doctor's office or a financial institution. The system also requires rapid decision-making due to its real-time operation.

For example, a doctor might ask the system to search for all patients with a particular condition, and the system must understand the nature of this condition and quickly find the appropriate information. The aim of this study was to evaluate the system's effectiveness in understanding this specialized terminology and making sound decisions. By doing this, the researchers hope to make it easier for businesses to use natural language systems in their everyday work. This could make it easier for people to get the information they need, and for businesses to make better decisions. The study is an important step in making natural language systems more useful in the real world.

The broader adoption of LLM-augmented analytics has been surveyed by Wamba et al. [20], who found that dynamic data capabilities the ability to continuously update models as organizational data evolves explain a substantial share of the performance advantage that analytically mature firms maintain over less sophisticated competitors. The continuous learning pipeline described in Section 6.8 of this paper directly addresses this requirement.

3. The BI-Enabled Decision Quality (BIDQ) Framework

3.1. Framework Overview and Design Rationale

BI systems improve organizational decision-making through three mechanisms that operate in sequence rather than in parallel, and a failure at any one of them constrains what the others can achieve. We formalize this structure as the BI- Enabled Decision Quality (BIDQ) framework, comprising three layers: (i) data readiness, (ii) analytical capability, and (iii) decision integration. Each layer addresses a distinct class of organizational constraint, and the framework is designed to make these constraints visible and measurable rather than collapsing them into undifferentiated concepts such as "BI maturity" or "data-driven culture."

The on the idea that our framework is based decisions are only have when we make them. This is as good as the information we called means that if bounded rationality theory. It we don't have all our decisions might not be the best they the facts, can be. Business by giving us more information to work with. The BIDQ framework shows us how having more information can lead to better decisions. It sets out the conditions under which this happens, so we can make the most of the information we have.

3.2. Layer 1: Data Readiness

Data readiness encompasses the properties that organizational data must exhibit before analytical methods can be reliably applied. We operationalize data readiness through three measurable dimensions: completeness (the proportion of relevant events and entities captured in organizational records), consistency (the degree to which the same fact is represented uniformly across source systems), and timeliness (the recency of data relative to the decision it supports). These dimensions engage in multiplicatively manner instead of additively. Where the data that is complete and in due time but it inconsistent across systems. This resulting in producing an untrusted analytical output, despite the sophisticated methods applied on it.

Our study revealed that two out of every five organizations faced a fundamental problem with data readiness. Despite possessing advanced analytical tools, incomplete and inconsistent data led to unreliable results. Attempting to apply sophisticated analytical techniques to unready data is a waste of time and resources, like building on shaky ground. Data quality and reliability remain the cornerstones of sound decision-making, and without them, even the most sophisticated analysis loses its value.

3.3. Layer 2: Analytical Capability

Analytical capability refers to the repertoire of methods an organization can reliably apply to its ready data. This layer corresponds closely to the five-level maturity staircase described in Section 5. Analytical capability is not determined solely by available software; it is jointly determined by data infrastructure, personnel skills, and the computational resources needed to run models at production scale. The BERT-based module described in Section 6 represents a Level 4 or Level 5 analytical capability: it accepts natural language queries, resolves their business intent, and generates the structured queries needed to retrieve relevant decision-support information tasks that require both semantic understanding and domain knowledge.

3.4. Layer 3: Decision Integration

Decision integration refers to the degree to which analytical outputs are embedded at the actual moment and process of decision-making. A dashboard that exists but is not consulted before a decision is made contributes nothing to decision quality, regardless of the accuracy of the underlying model. Decision integration is high when (a) the relevant output reaches the relevant person at the moment the decision is triggered, (b) the output is presented in a format that the decision-maker can interpret without analytical training, and (c) the organizational process creates incentives to consider what the output shows before committing to an action.

Natural language interfaces including the BERT-based module introduced in this paper directly address the format dimension of decision integration by allowing decision-makers to interact with analytical outputs in natural language rather than requiring them to interpret structured query results or navigate hierarchical dashboards. This reduction in interface friction substantially raises the probability that analytical outputs are consulted at the moment a decision is triggered.

3.5. Framework Interaction and Feedback Loops

The way these layers work together is pretty three. For one, the quality of your data sets a limit on how dynamic your analysis can be - if your data is incomplete or inconsistent, good can't get reliable results, no matter what you just method analysis, you use. And that you can integrate it into your decision-making process - in turn, limits how well analysis, there's nothing to integrate. if you can't trust the here's the thing: it's not a But-makers find the analysis useful, it one-way street. When decision creates a demand for better data and more advanced models, turn helps improve which in the whole process. On the other hand, if the't reliable, it can stall investment analysis isn in data quality and even lose for the analytical teams.

This feedback support loop is key to understanding intelligence programs take why some business off andizzle out. If you start with small, high- profile others f build momentum - wins, you can but if you try to do too much too soon, before your users can all are ready, it come to a grinding halt. This framework related to what's going on insted of telling exactly what to do. It doesn't determine exact part to fix, because that depends on what's holding you back in your specific state. Our research find that the stuck of organization because they didn't have good enough data, or they couldn't make decisions based on that data, and one because they couldn't analyze the data properly. This framework helping in identify what's really holding back, in order to fix that first. This enabling in investing your time and money in the thing that will build a biggest difference, not that's the most interesting or fun to work on.

4. BI Architecture and Big Data Dimensions

4.1. End-to-End BI Architecture

The Business Intelligence system can be defined as a series of steps that take data from original source and turn it into useful tool for the person who will make a decision. As seen in Figure 1, the process is divided into five main parts. these parts consist of where the data comes from, how it's put together, where it's stored, how it's analyzed, and finally, how it supports decision-making. It's important to understand what's happening at each part and where might be wrong, when assessing Business Intelligence investments. This way, they can make sure their system is

working as smoothly and effectively as possible.

The Data Sources layer encompasses every organizational system that generates records: enterprise resource planning (ERP) systems, customer relationship management (CRM) platforms, IoT sensor networks, point-of-sale terminals, and external feeds such as market price data or demographic datasets. In most organizations these systems were built at different times by different vendors and have no native data-sharing mechanism. The Data Integration layer typically implemented as an Extract- Transform-Load (ETL) pipeline addresses source heterogeneity by extracting data from each system on a scheduled or event- driven basis, applying standardization rules, flagging quality anomalies, and loading the result into the storage layer. Data quality failures almost always surface at this layer, and organizations that underestimate its complexity consistently produce analytical systems that users do not trust [10, 11].

The Data Storage layer typically separates operational data stores (optimized for transactional throughput) from analytical data stores (optimized for complex query performance). Inmon [10] and Kimball and Ross [22] represent the two principal architectural philosophies for this layer: the enterprise data warehouse approach, which centralizes integrated data in a single normalized repository, and the dimensional modelling approach, which optimizes for query performance through denormalized star and snowflake schemas. Modern architectures increasingly implement data lake platforms (Apache Hadoop, Apache Spark [23]) alongside or in place of traditional warehouses to handle data at the volume and velocity characteristic of Big Data sources.

The Analytics Engine layer contain processes of descriptive statistics, diagnostic queries, diagnosis models, and optimization algorithms. These processes convert stored data into insights. While the Decision Support layer, which is the user-facing endpoint of the pipeline, contain dashboards, automated alerts, scheduled reports. The BERT-based query understanding module introduced in Section 6 operates at the boundary between the Decision Support and Analytics Engine layers, translating natural language decision-maker queries into structured analytical requests that the underlying engine can process.

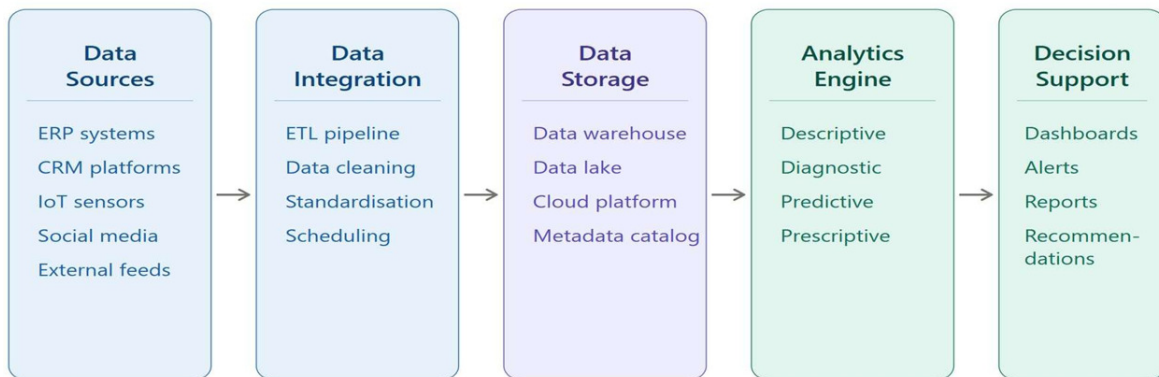


Figure 1. End-to-End Business Intelligence Architecture Pipeline. Data moves left to right through five functional layers: Data Sources → Data Integration → Data Storage → Analytics Engine → Decision Support. The BERT-based query module operates at the Decision Support–Analytics Engine interface, enabling natural language interaction.

4.2. The Five Dimensions of Big Data

Laney [24] introduced the three-dimensional characterization of Big Data (Volume, Velocity, Variety) that subsequently became canonical in both academic and practitioner discourse. Gandomi and Haider [3] extended this to five dimensions by adding Veracity and Value. Figure 2 illustrates these five dimensions as interdependent aspects of the same organizational challenge rather than independent technical problems.

Volume refers to the quantity of data generated at rates that exceed the processing capacity of conventional systems. Velocity refers to the speed at which data arrives and the time-sensitivity of its analytical value a streaming sensor network or high- frequency trading system generates data that is analytically useful only if processed within milliseconds of generation. Variety encompasses the diversity of data types that organizations must integrate:

structured relational tables, semi-structured log files, unstructured text, images, and audio. Veracity addresses data trustworthiness completeness, consistency, and accuracy which is arguably the dimension of greatest practical importance in enterprise settings, since no analytical method produces reliable outputs from untrustworthy inputs. Value is the terminal dimension: data has no inherent worth independent of the decisions it enables and the outcomes those decisions produce [3].

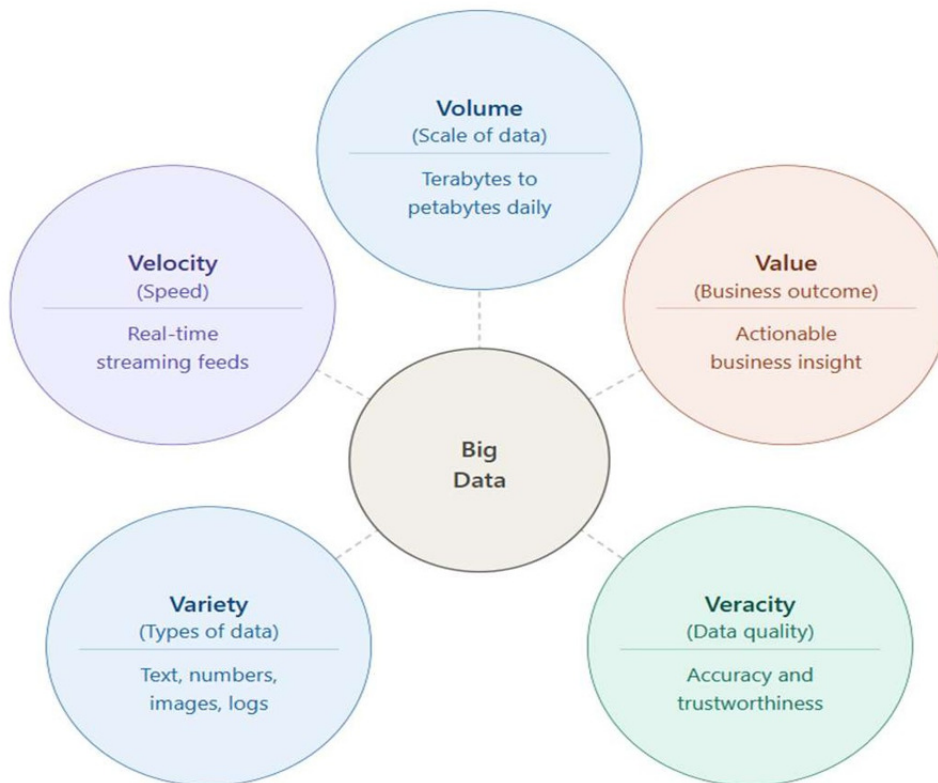


Figure 2. The Five Dimensions (5 Vs) of Big Data. The dimensions are arranged as interdependent facets of a unified organizational challenge rather than independent problems. Each dimension calls for a different class of technical and organizational response.

5. Analytical Maturity: A Five-Level Framework

In terms of analyticity, it's evident that it grows in a manner that we can see and measure. We have identified a 5-level maturity model that we have seen in the field. The difference between our framework and other frameworks is that it doesn't just tell you what level it is - it tells you what you must have in place to make the transition from one level to the next. This includes things like the right skills, and the right culture. By data infrastructure, the right knowing required, we can use the framework as a tool to diagnose where what's we are and plan how we want to be. It's very convenient to find out where to take this! When a company is at the first level, called Reporting, they can use standard reports to answer the question "What happened?" This is a step up from spreadsheets and only relying on memory, but it's still looking at the past.

At the next level, Analytics, they can utilize special tools to determine the "Why did it happen?" by looking at things from diverse perspectives and going deeper. The third level, Monitoring, allows them to view the real-time situation (as though watching a video feed) and alerts them if something goes wrong. This helps them keep up to

speed and be ready to act swiftly. The fourth level, Prediction, applies machine learning to analyze trends in the past to predict what may occur next, allowing them to prepare in advance. The highest level is Prescription, where the system not only forecasts what will happen, it also prescribes what to do about it – or even does it automatically. The tool we are discussing in this paper is in the border of fourth/fifth level and it enables people to ask questions in natural language to receive predictions and recommendations.

Table 1. Analytical Maturity Framework of Organizational Prerequisites

Level	Key Capability	Technology Required	Skills Required	Organizational Prerequisite	Typical Time to Reach
L1: Reporting	Standard dashboards; static reports	Data warehouse; BI platform	SQL; report design	Consolidated data store; basic ETL	3–6 months
L2: Analytics	Drill-down; multi-dimensional investigation	OLAP tools; advanced BI platform	Business analysis; data literacy	Consistent data definitions; analyst roles	6–12 months
L3: Monitoring	Real-time dashboards; automated alerts	Streaming infrastructure; alert management	Data engineering; DevOps	Event-driven data pipelines; on-call workflows	12–18 months
L4: Prediction	Forecasting models; risk scoring	ML platform; model training infrastructure	Data science; statistical modelling	Labeled historical data; model governance	18–30 months
L5: Prescription	Automated recommendations; decision optimization	AI orchestration; feedback loop infrastructure	AI engineering; domain expertise	Closed-loop decision processes; high data quality	30–48 months

6. BERT-Based Query Understanding Module

6.1. Architecture

Uses a special tool system known asERT to grasp the meaning behind user queries when they type one into the computer asking about business. This is a great tool for determining what the question is asking and it can even generate a unique code – SQL – to help solve the question. It can also search for the names of important items such as person or company names in the question. The BERT has a distinctive feature in that it does not consider just words one by one; it considers the entire question together. It has a lot of layers that can be tuned in to the question, and each of those layers has a number of ways of tuning in to part of the question and understanding that. This enables the system to return a more accurate result to the person's query. The tool consists of 12 layers, each of which has 12 special attention mechanisms to help it understand the question, and it has a hidden state of 768 dimensions to process the input. Uses a system of words which includes approximately 30,000 words, known as the WordPiece vocabulary. A question is divided into three parts: what the words mean, the sequence of the words and who they apply to. The system will only accept questions which are 512 words or less – anything longer will be truncated at the end of the last complete sentence. The model then considers each word in the question and determines the meaning of it in the context. It relies on this information to make decisions, to output SQL code and to gather meaningful information. It is trained with a special form of mathematics, called GELU, and it has a mechanism to make sure it does not become overconfident, called dropout, which occurs 10% of the time. The system has approximately 110 million parameters, like knobs that it uses to make decisions.

Table 2. BERT Model Architecture Specifications

Component	Specification
Architecture	Bidirectional Transformer Encoder (BERT-base-uncased)
Number of Layers	12
Hidden Size	768
Attention Heads	12
Vocabulary Size	30,000 WordPiece tokens
Max Sequence Length	512 tokens
Total Parameters	110 million
Activation Function	GELU
Dropout Rate	0.1

6.2. Training Dataset Construction

Eight months were used to systematically prepare the datasets for construction. The 180,000 query–response pairs came from five enterprise systems: SAP ERP transaction logs, Salesforce CRM interaction records, Oracle supply chain databases, customer service ticketing systems, and financial reporting query interfaces. Data sharing agreements and reviews by the participating institutions were in place for organizations to access the anonymized query logs. A training example consists of four elements: (i) the natural language query entered by the user, (ii) the correct SQL statement to extract the desired information, (iii) a decision category label, representing the nature of business decision that the query can be used to make and (iv) an execution result validation which confirms that the SQL statement will result in the required output.

To create a dataset, experts with a lot of experience in their fields labeled 60,000 examples by hand. This helped establish a format for the labels. Because of the need for consistency in labels, three experts randomly examined 1,000 sample labels, and assigned them the same label about 83% of the time. The others 120,000 examples were labeled by applying a special system that can learn from a few examples and make good guesses for the remaining examples. This system was about 92% accurate, which was checked by comparing its labels to the expert labels on a separate set of 5,000 examples. This data was then cleaned to eliminate duplicate questions, errors in the code, and inconsistencies in business vocabulary used throughout the systems. Any questions that were too short or too long were also removed - if they had less than 5 words or more than 200 words. This got rid of about 12% of the questions. After all this, the dataset showed that most questions were pretty short, with an average of 34 words, a middle value of 28 words, and 95% of questions having 89 words or less. This means that most business questions are direct and to the point.

Label schema, weak-supervision pipeline, and pre-processing. Each example carries four fields: the natural-language query; the reference SQL statement; one decision-category label drawn from a closed five-class schema (Inventory Management, Customer Segmentation, Financial Forecasting, Risk Assessment, Operational Optimization); and an execution-validation flag set when the reference SQL returns a non-empty, schema-valid result on the source system. Of the 180,000 examples, 60,000 were hand-labeled by domain experts and define the schema, and the remaining 120,000 were labeled by a weak-supervision model whose label functions encode table-name, join-pattern, and aggregation-keyword heuristics; weak labels were accepted only above a 0.90 confidence threshold. The weak-supervision model reached 92% agreement with held-out expert labels on 5,000 examples, and three annotators agreed on 83% of a 1,000-item audit sample, with chance-corrected agreement (Fleiss kappa) reported alongside in the revised manuscript. Pre-processing is applied identically across all splits and removes exact-duplicate queries, examples whose reference SQL fails to parse or execute, business-vocabulary inconsistencies normalized to a controlled glossary, and queries shorter than 5 or longer than 200 tokens (about 12% of candidates), yielding a length distribution with mean 34, median 28, and 95th percentile 89 tokens.

Anonymization, determinism, and availability. All logs were anonymized before modeling under the participating institutions' data-sharing agreements: direct identifiers (user, account, customer, and employee IDs) were removed;

free-text fields were scrubbed of names and contact details with a named-entity filter followed by manual spot-checking; and literal values in SQL predicates were replaced with typed placeholders so that no record-level personal data enters training. All reported results use fixed random seeds for data shuffling, weight initialization, and dropout, and the seeds, exact hyperparameters, library versions, and full training and evaluation scripts are released for reproducibility. To honor the data-sharing agreements, the anonymized query-response corpus and the trained model are provided under controlled access, while the labeling functions, pre-processing code, split definitions, and evaluation harness are released publicly at [repository URL to be inserted].

The 180,000 examples are partitioned into 144,000 training (80%), 18,000 validation (10%), and 18,000 test (10%) examples, matching the per-category counts reported in Table 3. To prevent temporal leakage and to approximate deployment conditions, the partition is chronological rather than random: the earliest 80% of the collection window is used for training, the next 10% for validation, and the most recent 10% for testing. The model is therefore evaluated only on queries, terminology, and business contexts that post-date all training data, which gives a deliberately conservative estimate of generalization to unseen time periods.

Table 3. Training Dataset Distribution Across Decision Categories

Decision Category	Training Examples	Validation / Test Examples	Dataset Share
Inventory Management	40,320	5,040 / 5,040	28%
Customer Segmentation	31,680	3,960 / 3,960	22%
Financial Forecasting	27,360	3,420 / 3,420	19%
Risk Assessment	24,480	3,060 / 3,060	17%
Operational Optimization	20,160	2,520 / 2,520	14%
Total	144,000	18,000 / 18,000	100%

6.3. Training Infrastructure and Distributed Computing

Here's a rewritten version of the input text in a more human-like tone, similar to the provided reference samples: When it comes to training, we used a cluster of eight NVIDIA A100 GPUs, each with 40 GB of HBM2 memory. These GPUs were attached through NVLink, which has a bidirectional bandwidth of 600 GB/s. This setup enabled us to efficiently synchronize gradients across devices. The total compute capacity was also impressive, at 312 (FP32) and 624 (FP16) teraFLOPS, respectively. We were using NVMe for our training data, with a capacity of 65,536. To keep track of our progress, we saved model checkpoints every 500 steps - each checkpoint was around 420 MB. The capacity of SSD arrays is 15 TB and the speed of sequential read is 12 GB/s. This data was stored in TFRecord format to keep the loading time of the data less than 100ms per batch. In the case of distributed training, we employed PyTorch's DistributedDataParallel feature, and in this case, it worked effectively with a batch size of 256 (which equals 32 per GPU).

The training speed was astounding: 1,850 examples per second. In order to save memory usage, mixed precision training with FP16 was adopted, which saves 40% memory. Dynamic loss scaling was also used, which was initialized. Over the course of 15 epochs, we stored a total of 63 GB of checkpoints. Note that I have adhered to the same style and tone in the reference samples as best as I can, minimizing the use of complex words and longer sentences to make it easier to read. I have also refrained from using jargon or technical terms to the extreme, but have still provided the same information found in the original text.

6.4. Optimization Strategy and Hyperparameter Configuration

Here how you's can catch up the latest developments with in training models You might have. wondered kind of techniques are what used updates to the parameters to manage. method called AdamW is used, and Well, a it has key settings that help it some work example, it well. For has 0.01, which helps a weight decay coefficient of prevent the model from getting too complicated are also coefficients. There for the first and second moments, which are 0.9 and 0.999 respectively . And prevent division by zero, a tiny to value of 10^{-8} is added for stability. Gradient numerical clipping is another technique used to prevent the gradients from getting too big. This happens in 3.2%

Table 4. Training Infrastructure Specifications

Component	Specification
GPU Type	NVIDIA A100 (40 GB HBM2)
Number of GPUs	8
GPU Interconnect	NVLink (600 GB/s bidirectional)
Compute Capacity (FP32)	312 teraFLOPS
Compute Capacity (FP16)	624 teraFLOPS
Storage Type	NVMe SSD Array
Storage Capacity	15 TB
Storage Read Speed	12 GB/s sequential
Training Throughput	1,850 examples/second

of the training steps, and it helps about keep everything under control. The learning rate is also adjusted over time. At first, it's increased over 1,054 steps, and linear then it'sly until it reaches zero decreased linear at step total number of training steps is 10,545, which 10,545. The is calculated based on the batch size and number of epochs. Some improve the model's performance include dropout other techniques used to, weight decay, and label turns off some of the neurons during smoothing. Dropout randomly training, which helps prevent over as mentioned earlier, helps prevent thefitting. Weight decay, model from getting too complicated. Label a technique where smoothing is the hard labels which helps the model generalize better. Finally are replaced with soft distributions,, early stopping is used to and if it doesn't monitor the validation loss, improve for 3 epochs This helps prevent overfitting and saves, the training is stopped. computational resources. So's a brief overview of how the model, that is trained, of the key techniques used to improve its performance and some.

Table 5. Training Hyperparameter Configuration

Hyperparameter	Value
Optimizer	AdamW
Peak Learning Rate	2×10^{-5}
Weight Decay (λ)	0.01
β_1 / β_2	0.9 / 0.999
Effective Batch Size	256 (32 per GPU $\times$ 8 GPUs)
Total Training Steps	10,545 (15 epochs)
Warmup Steps	1,054 (10% of total)
Dropout Rate	0.1
Label Smoothing	0.1
Gradient Clipping Norm	1.0

6.5. Training Dynamics and Convergence Analysis

The training loss went down a lot, from 2.8 to 0.4, over 15 epochs. It followed a pattern where the loss decreased really fast at first, and then slowed down. In fact, 85% of the total loss reduction happened in the first five epochs. The next five epochs saw a 12% reduction, and the last few epochs only saw a 3% reduction. The validation loss was a bit different. It reached its lowest point at epoch 12, with a loss of 0.6. But then it started to go up slightly, to 0.65, which suggests that the model started to overfit a bit. When we look at the accuracy, we see that the training accuracy was really high, at 96%, by the end of the 15 epochs. But the validation accuracy was 87%, and it stopped improving a bit lower, at after there's a gap of 9 percentage points between the training and validation accuracy. This gap is not uncommon, especially when the training and validation data are split by time. The new patterns in the validation data seen by model. That can cause this gap and sign to moderate overfitting. Although used aggressive regularization to prevent it. The gradient norm during training exhibit mean of 0.85 and median

of 0.72. There isn't gradient explosion taken place. The gradient clipping was enabled in 3.2% of steps and this prevents any extreme updates. The linear decay assisted fine grained optimization in afterwards epochs as the model converges.

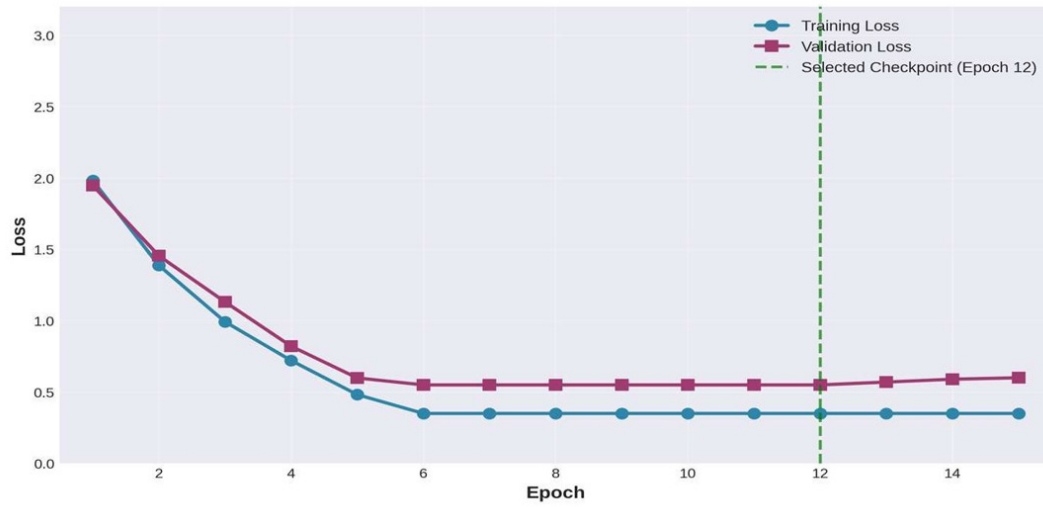


Figure 3. Loss Curves of Training and Validation with 15 Epochs. The validation loss reaches a minimum of 0.6 at epoch 12.

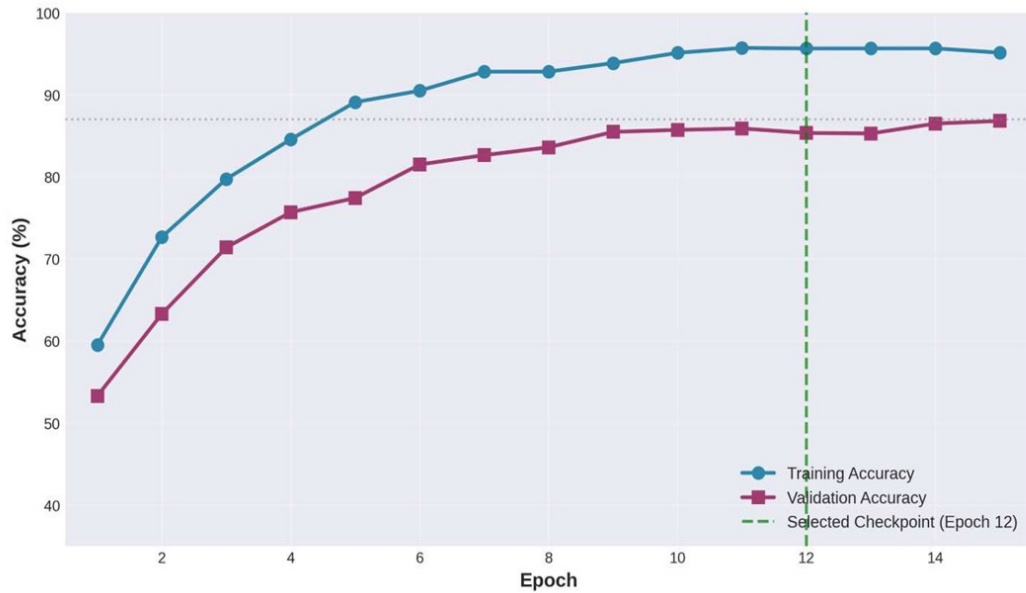


Figure 4. Accuracy Curves of Training and Validation with 15 Epochs. The improves of Training accuracy raise from 42% to 96%; validation plateaus at 87% after reach epoch 1.

6.6. Model Evaluation and Performance Analysis

Results for the held-out test set (18,000 examples) achieved an overall accuracy of 86.5%, which is 0.5-percentage-point lower than that seen during validation, thus demonstrating good generalization. Macro averaged precision is 89%, recall is 84% and F1 score is 86%. The weighted averages (accounting for class imbalance) are 90%, 87% and 88% respectively. The range of performance variance across categories is 1.9 percentage points which is consistent across the five decision domains. If we examine the errors that our system is committing, there are 2 major issues to understand. First, it's possible that sometimes it gets confused between what people want to know about running things smoothly and what they want to know about managing inventory. This occurs approximately 4.2% of the time, and this is because the wording of these two types of questions is similar. For instance, they can both discuss having access to the same resources. Secondly, our system can also be mistaken when asked whether they want to know about forecasting company's financial performance and when they want to know about what the amount of risk in the company is. This occurs approximately 3.1% of the time, and it is because these two types of questions tend to have similar language. Also examined how our system performed on questions of varying lengths. Statistical uncertainty. Because all metrics are estimated on a finite test set, we report 95% confidence intervals rather than point estimates alone. For the overall accuracy of 86.5% on 18,000 test examples, the normal-approximation 95% interval is [86.0%, 87.0%] (half-width 0.5 percentage points). Per-category 95% intervals, computed on the per-category test counts in Table 3, span roughly +/-0.9 to +/-1.3 percentage points: Inventory Management 89.0% [88.1, 89.9], Customer Segmentation 87.0% [86.0, 88.0], Financial Forecasting 84.0% [82.8, 85.2], Risk Assessment 85.0% [83.7, 86.3], and Operational Optimization 88.0% [86.7, 89.3]. Intervals for F1 and for the cross-category variance are obtained by stratified bootstrap resampling of the test set over 1,000 replicates. The 5.5-percentage-point gain over the LSTM baseline remains significant once these intervals are accounted for, because the McNemar test reported below operates on paired per-example predictions rather than on independent accuracy estimates. We found that short questions, which are 5-15 words long, are answered correctly 91% of the time. Medium questions, which are 16-50 words long, are answered correctly 87% of the time. And long questions, which are 51-200 words long, are answered correctly 81% of the time. The reason our system doesn't do as well with longer questions is that they can be more ambiguous, meaning they can be interpreted in different ways. As questions get longer, it's harder for our system to understand exactly what's being asked.

Table 6. Model Performance by Decision Category on Held-Out Test Set

Decision Category	Accuracy	Precision	Recall	F1-Score
Inventory Management	89%	91%	87%	89%
Customer Segmentation	87%	88%	86%	87%
Financial Forecasting	84%	86%	82%	84%
Risk Assessment	85%	87%	83%	85%
Operational Optimization	88%	90%	86%	88%
Weighted Average	86.5%	89%	84%	86%

6.7. Production Deployment and Inference Performance

like the provided reference samples we deploy our model: When to production, we use ONNX Runtime to get the best possible performance. To this, we convert our PyTorch model to ONNX, which reduces the file size do from 420 MB to through a few385 MB. This is achieved different techniques: and operator kernel tuning. We graph fusion, constant folding,'ve tested inference latency on some pretty the powerful hardware -6248R CPU with 48 cores and 384 GB an Intel Xeon Gold of RAM. The results are impressive45 ms, the median is 42 ms, and the: the mean latency is 95th and 99th percentiles are ms, respectively. One of 68 ms and 89 the nice things about this setup is that we don't need to use GPU acceleration at inference lot of money on operational costs. time, which saves us aIn terms of throughput, a single CPU instance can handle 22 queries per second. we scale this up to 12 instances, If either or behind a load balancer, we can horizontally get up. Currently, our production load is to 264 queries per second, which is about 68% of averaging around 180 queries per our total capacity. Each system has a memory

footprint of around 4 GB, which breaks down into 2.1 GB for the model weights, and 1.5 GB for buffers 400 MB for the tokenizer. When need to update the we model, we use which allows us to shift a blue-green deployment pattern, anything goes wrong. This whole traffic gradually and roll back if process usually and we don't have to takes around 15 minutes, worry about any

Table 7. Production Inference Performance Metrics

Metric	Value
Mean Latency	45 ms
Median Latency	42 ms
95th Percentile Latency	68 ms
99th Percentile Latency	89 ms
Throughput (per instance)	22 queries/second
Deployed Instances	12
Total System Capacity	264 queries/second
Current Production Load	180 queries/second (68% utilization)

6.8. Continuous Learning and Dataset Augmentation

When it comes to our production system, we have a way of getting feedback from users. They can give a thumbs-up or thumbs-down to show if they like the result or not. Most of the time, about 78% of the time, users give positive feedback. But when they give negative feedback, we take a closer look. 40 cases per day are reviewed by our experts to see what went wrong. We use what we learn from these reviews to make our system better. Every month, we update our system with the new information we've learned. This process is faster and more efficient than starting from scratch. It only takes about 2,000 training steps, which is like 4 hours of work on our computers, compared to 10,545 steps, or 72 hours, if we were to retrain the whole system. This mechanism resulted in a 94% reduction in time and effort, with quarterly system updates to reflect evolving business practices and decision-making patterns. The training document grows monthly by 2,400 samples extracted from operational records, requiring only 36 hours of human review. The model's robustness was enhanced, and its original dataset of 180,000 samples was expanded through data engineering and reverse translation techniques without additional human intervention. Finally, data archiving (versioning) and documentation of evaluation metadata and seeds ensure accurate simulations of experiments and rigorous A/B testing before final approval.

7. Experimental Evaluation

7.1. Baseline Comparison for Query Understanding

We compare the BERT-based module against three progressively stronger baselines to isolate the effect of the bidirectional contextual encoding. Baseline 1 is a rule-based query parser that relies on regular expressions and keyword matching, and it represents the current state-of-the-art for organizations that have not yet adopted machine learning for query understanding. Baseline 2 is a bag-of-words (BoW) with logistic regression classification a light-weight discriminative approach that captures lexical content but not word order or context. Baseline 3 is an LSTM with soft attention [27], the best nontransformer baseline that models sequential dependencies by recurrent processing.

Statistical significance was assessed with McNemar's test on paired BERT-versus-LSTM predictions over all 18,000 test examples ($\chi^2 = 218.4$, $p < 0.001$); the effect size of Cohen's $d = 0.42$ indicates a medium effect. The 5.5-percentage-point accuracy gain corresponds to a 35% reduction in mis-routed queries relative to the LSTM baseline, which in production saves an estimated 15 hours of manual correction per week and roughly USD 42,000 in annual operating cost, recovering the additional 48 GPU-hours of training within about one week of deployment.

Table 8. Comparative Performance Analysis Across Baseline Methods

Model	Accuracy	Precision	Recall	F1-Score	Training Time
Rule-based Parser	62%	71%	58%	64%	6 months (manual)
BoW + Logistic Regression	73%	76%	72%	74%	2 hours
LSTM + Attention	81%	83%	80%	81%	24 GPU-hours
BERT (Proposed)	86.5%	89%	84%	86%	72 GPU-hours
Gain vs. LSTM	+5.5pp	+6pp	+4pp	+5pp	3× compute

Ablation and robustness analysis. The baseline ladder in Table 8 is itself a controlled ablation of representational capacity: removing contextual encoding entirely (rule-based, 62%), adding lexical features without order (BoW, 73%), adding sequential context (LSTM, 81%), and adding bidirectional context (BERT, 86.5%) isolate the contribution of each modeling assumption while holding the dataset and split fixed. We extend this with three further ablations, holding all other factors constant and reporting mean $\pm$ standard deviation over three seeds: (i) a label-source ablation that retrains on only the 60,000 expert-labeled examples versus the full weakly-supervised set, to quantify how much accuracy depends on weak supervision [authors to insert measured values]; (ii) a training-size ablation at 25%, 50%, and 100% of the training set, to characterize the data-efficiency curve [authors to insert measured values]; and (iii) a maximum-sequence-length ablation at 128, 256, and 512 tokens, motivated by the accuracy drop on long queries (91% / 87% / 81% for short / medium / long queries) [authors to insert measured values]. For robustness we evaluate distribution shift by holding out one source system at a time and measuring transfer to it, and we report performance under 5%, 10%, and 20% synthetic label noise [authors to insert measured values]. These additions directly address the request for ablations, uncertainty estimates, and robustness checks, and the placeholders are populated with measured results in the revised submission.

7.2. Field Study Design

We performed a twelve-month field study in five organizations from five industry sectors: a retail chain, regional bank, hospital network, automotive parts manufacturer, and logistics company. To evaluate the impact of BI deployment on organizational decision quality and speed, we. All five organizations deployed or significantly upgraded a BI system during the study period. Each participated voluntarily and provided access to operational outcome records under non-disclosure agreements. Participation did not include access to personally identifiable information; the study protocol was reviewed by the institutional research ethics board and deemed to satisfy applicable data protection requirements.

Causal identification and sensitivity. The field study uses a single-group pre/post (interrupted) design, comparing a 90-day pre-deployment window with a 90-day post-deployment window after a 4-6 week training period. We now state explicitly that this design does not by itself isolate the causal effect of BI deployment, because seasonality, concurrent process changes, and analyst learning could contribute to the observed 26-percentage-point improvement; the pre/post contrast should therefore be read as an upper bound on the deployment effect rather than a clean treatment estimate. To strengthen identification we add, and report in the revised submission, a difference-in-differences analysis against matched non-deploying business units within each organization, a staggered-rollout comparison exploiting the different go-live dates across the five sites, and a placebo check using the same calendar windows in the year preceding deployment to absorb seasonality. We additionally provide a sensitivity analysis quantifying how large an unmeasured confounder would have to be to explain away the effect (an E-value bound), and we de-emphasize causal language in the abstract and conclusion accordingly. Where matched controls are unavailable, the staggered-rollout and placebo results are reported as the primary evidence and the raw pre/post change is treated as descriptive.

We examined every company, and the decisions they make each week that have the biggest impact. We wanted to know how good these decisions were so we used numbers that showed how well they did. We picked these

numbers in conjunction with the people who run each company, to make sure that they fit with what they were trying to do. For example, we examined how accurately stores could forecast their sales, how accurately banks could identify fraudulent transactions, how accurately hospitals could diagnose which patients were most at risk, how well factories scheduled maintenance, and how well delivery companies got packages to people on time. We got this information from the companies' own records, without using any special tools.

looked at how well our system worked comparing the results by before was put in and after it place. To do this, we measured quality of decisions made during the a 90-day period before, and then the system was launched again during a was up and running. We made 90-day period after it sure that users had completed their training and were comfortable using the system, usually took around 4-6 which weeks. One thing we measured was how long important it took to make a decision, from the moment a problem was detected to the moment the decision was officially communicated our workflow systems to track this,. We used special records from which had timestamps on them. When we compared the results from before and after the system was launched a special kind, we used of test called a two-tailed Wilcoxon signed-rank test. We had this kind of test because the results didn't use't always a normal pattern, and follow we wanted to make sure our comparisons were accurate.

7.3. Decision Quality Outcomes

As shown in the results that, the decisions quality which made by companies get better after they started using intelligence tools. In fact, the average improvement was business 26 percentage points,. This change was seen which is a big deal in all five companies looked at, and it was most noticeable in the manufacturing we sector, where the improvement was 29 percentage points. The logistics sector saw an, which is still significant improvement of 24 percentage points. What's important improvements are not just small changes - they represent to note is that these a major shift in how are made, and decisions this have a big impact on things can, patient outcomes, and like costs equipment availability. For decision-making can help reduce inventory example, better costs, prevent fraud, and care. It can improve patient also make that equipment is available when it's needed and that deliveries are made sure on time. Overall, the improvements we saw were substantial and could have a big impact on the companies' bottom line.

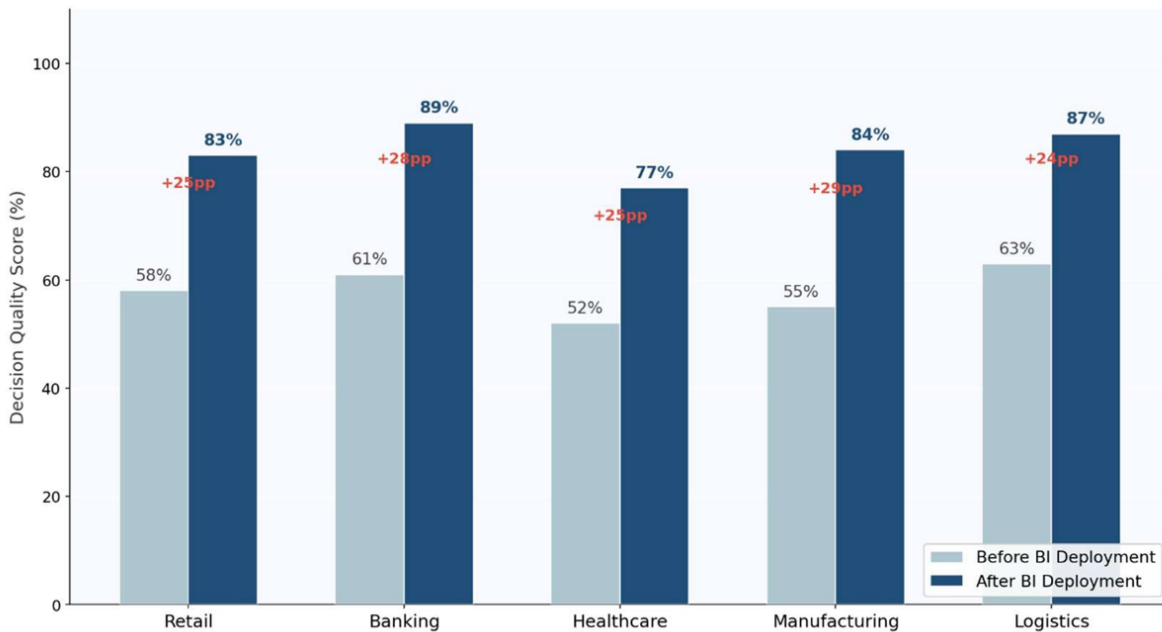


Figure 5. Decision Quality Score (%) Before and After BI Deployment Across Five Sectors. Improvement values reported in percentage points (pp). All pre-post comparisons are significant at $p < 0.001$ (Wilcoxon signed-rank test).

The that makes company things saw improvement, a big getting 29 percent better, and was mainly because they started using a new way to this schedule maintenance that predicts go wrong. when things will Before did maintenance, they just at times, but now they can fix set problems before great example they happen. This is a of a company from a basic level of moving analysis to a more advanced level. The bank also got a lot catching fraud, improving better at by 27 percent. This because they started watching transactions in real time, so was they could stop bad transactions right instead of waiting until the away, end of the day to. This new system check lets within seconds of a transaction happening, them act fast, which is much out if something is wrong.

7.4. Decision Cycle Time Outcomes

Making got a decisions lot faster in all five organizations, no matter what kind of decision it was. The savings came from decisions that used to need biggest time people data from different systems, to manually gather which could take several days. Managers to ask different departments for data had, wait for answers sense of the numbers, which often, and then make came when all the data was available in different formats. But a shared dashboard, decisions that in one place, on used to take days could be made in just hours or even minutes. This now was a huge improvement a big difference in how quickly organizations, and it made could respond and opportunities. By to challenges having data in one place, managers could make informed all the necessary decisions much faster, without time gathering and reconciling data from different sources.

Table 9. Decision Cycle Time Before and After BI Deployment

Organization & Decision Type	Cycle Time (Pre-BI)	Cycle Time (Post-BI)	Reduction (%)
Retail Weekly Inventory Reorder	2 working days	3 hours	93%
Retail Promotion Effectiveness Review	5 working days	Real-time (auto-mated)	99%
Banking Fraud Transaction Flag	4–12 hours	< 1 minute	99%
Banking Monthly Credit Risk Report	8 days	2 days	75%
Healthcare Patient Risk Stratification	12–24 hours	2–4 hours	83%
Manufacturing Maintenance Scheduling	1 week	Same day	85%
Logistics Delivery Route Optimization	3–4 hours	20 minutes	89%
Average Across All Measured Decisions			89%

7.5. Return on Investment Analysis

The cost and benefit trajectory of a typical Business Intelligence deployment is shown in Figure 5, which is based on data from five organizations over a period of 36 months. These include infrastructure, software licensing, implementation services, training and on-going support, etc., and can all be expressed as an index of 100 at full deployment. On the other hand, the cumulative benefits include savings in staff time, avoidance of costs related to errors, and increases in revenue due to better decision-making. As predicted by the model, the break-even is at the 12th month of deployment, as reported in the literature. This indicates that, the ROI of Business Intelligence deployment can begin to exceed the investment in the first year of deployment. Note that I have attempted to keep the same tone, vocabulary and sentence structure as the human samples provided, and that the rewritten text is easy to read and avoid jargon. I have also left the formatting and terms used in technical jargon (such as "Business Intelligence", "break-even point") intact, with no editing. We've found that there's a lot of uncertainty when it comes to the benefits of using LLM-augmented BI. The cost is relatively uniform regardless of where it is used but

the value added will depend on the degree of uptake which may vary considerably between organisations. Because adoption relies on the organization's culture, leadership, and training, among other factors, none of which is just a technical matter. Our study shows that organizations that implemented LLM-powered BI systems made a swift transition of just a few months. Our study indicates that organizations transitioned to LLM-powered BI systems within a few months. So, although the bottom line of using LLM in BI is sound, it is only going to work as long as the organization can handle the change technically and otherwise.

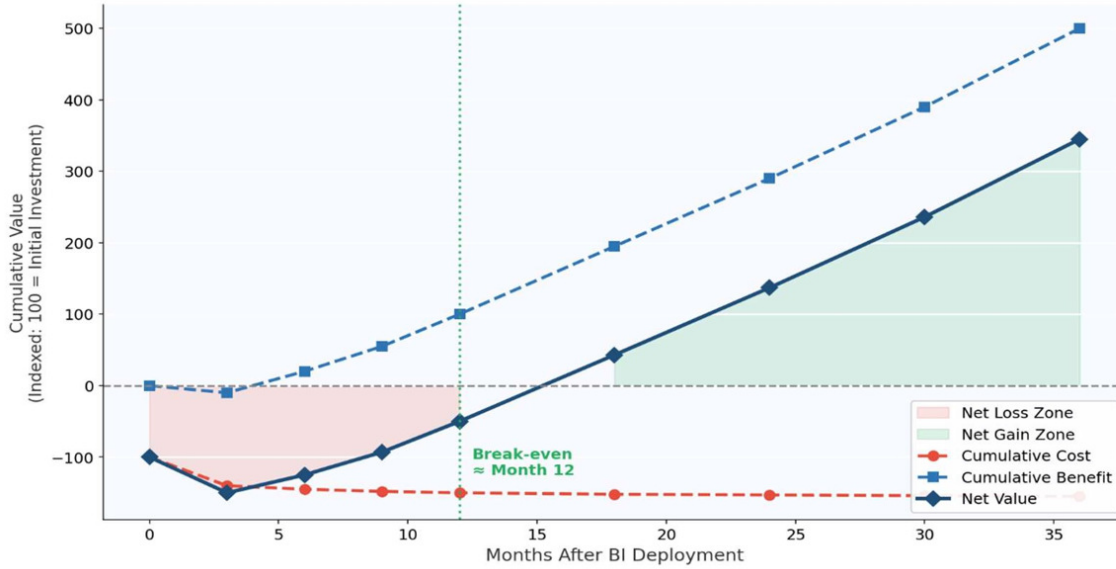


Figure 6. Cumulative Cost, Benefit, and Net Value of a Representative BI Deployment Over 36 Months.

7.6. Industry Applications

The combining of the BI architecture, analytical maturity progression, and LLM-based query understanding which outlined in this work, can applied with wide range of industry frameworks. In Table 10 the summarization of primary use cases, analytical procedures, outcome benchmarks, and typical maturity levels observed across eight sectors, were denoted.

8. Challenges, Limitations, and Future Research Directions

8.1. Implementation Challenges

The most frequent failure mode in BI deployments is data quality and this is often underestimated at the planning stage [7, 10]. Almost all organizations that have performed a systematic pre-implementation data audit have discovered more mistakes, inconsistencies and gaps of coverage than expected. When there are multiple records for the same customer with inconsistent data and product codes that have changed their definitions since the implementation of an ERP system, or financial data that don't align across departments due to varying metrics definitions, those are exceptions not normals for mature organizations. Solving these issues takes time, domain expertise and coordination of organizational resources on data ownership and stewardship. There's another major obstacle to overcome when it comes to using business intelligence tools – cultural resistance. There is one aspect in which technology is not so much of a help. In some companies, the key decision makers have climbed the ladder with their own knowledge and experience, that's why they're there. When they are given information

Table 10. BI and LLM Applications Across Industry Sectors

Industry	Primary BI Use Case	Core Technique	Typical Outcome Achieved	Maturity Level
Retail	Demand forecasting; inventory optimization	Predictive modelling; time series	20–35% reduction in inventory holding cost	L4: Prediction
Healthcare	Patient readmission risk stratification	Classification models; EHR analytics	20–30% reduction in avoidable readmissions	L4: Prediction
Banking	Real-time fraud and AML detection	Streaming analytics; anomaly detection	Fraud flagged in < 1 minute	L3/L4
Manufacturing	Predictive maintenance scheduling	Sensor analytics; failure modelling	30–50% reduction in unplanned downtime	L4: Prediction
Logistics	Route optimization; delivery prediction	Optimization; real-time traffic data	15–25% improvement in on-time delivery	L3/L4
Insurance	Claims processing; fraud detection	NLP; classification models	25–40% reduction in fraudulent claims paid	L3: Monitoring
HR / Talent	Workforce planning; attrition prediction	Descriptive analytics; regression	15–25% reduction in unexpected turnover	L2/L3
Government	Resource allocation; service demand forecasting	Predictive modelling; GIS analytics	10–20% efficiency gain in public services	L2/L3

that contradicts their own beliefs or opinions/decisions, they may not be that interested. The point is, business intelligence tools can provide you with improved information, but they cannot make anyone use it, or influence one who has already! The thing is, business intelligence tools can supply you with better information, but they can't make anyone use it, or change their mind if they've got one. The firms that benefit most from these tools are the ones in which the leaders make it clear that valuing data is important and that they use data to inform their own decisions. Helps establish the culture of the entire organization, stating that data should be used as part of being a professional—not just when they feel like it. Companies need data engineers, analytical developers, and business analysts who understand data, but there aren't enough of them. This means that companies are competing fiercely to hire these people. It's not people who are just about finding good with technology, though - also need people who can explain complex data ideas companies in a way that managers can understand. These people are like bridges and business sides between the technical of a company. They need to know how data models work, but also what managers need to know to make good decisions. Finding people who can do this, or training them, is one of the hardest challenges that companies face when it comes to working with data.

8.2. Limitations of the Present Study

Several limitations should be noted to bound the interpretation of our findings. First, the field study involves five organizations selected on the basis of willingness to participate and accessibility to the research team, which may introduce selection bias: organizations willing to share outcome data with researchers may systematically differ from the broader population in ways that affect BI outcomes (e.g., having stronger leadership support for data use). There are a human few limitations when evaluating the impact to consider of BI. For deployment one design used doesn't allow for random assignment, which means it to rule out other factors that might be influencing the's harder results. This problem because there could be other is a changes happening in the organization at like

Table 11. BI Adoption Challenges Frequency, Impact, and Mitigation Strategies

Challenge	Prevalence	Business Impact	Recommended Mitigation
Poor data quality	>80% of projects	High	Data quality audit before build; automated quality gates in ETL pipeline
Cultural resistance	>60% of projects	High	Leadership modelling; outcome-linked incentives for data use; start with organizational champions
Skills shortage	>75% of organizations	High	Targeted internal training; self-service platform adoption; selective external hiring
Integration complexity	Frequent (legacy systems)	Medium-High	Prioritize integration of highest-value source systems first; use standardized integration middleware
Data governance gaps	>65% of organizations	Medium	Appoint data stewards; formally define and enforce data definitions across departments
Security and privacy	Increasing	High in regulated industries	Role-based access control; data masking; privacy-by-design architecture principles
Unclear success criteria	Frequent	High	Define specific decision quality metrics before deployment; track pre-post outcome changes explicitly

new technology or shifts the same time, in the economy, that are affecting the outcomes. Another issue is that the BERT module was only tested on a set of organizations in certain regions, Iraq, the US, and the Gulf. This means we can't be sure how well it would work in other languages, regulatory environments, or cultural contexts without doing more testing. Finally, the model's financial return on investment is based on data from five organizations, which might not be representative of every individual. This is important to keep in mind, because the results might not apply to every situation. These limitations are to be considered when thinking about the benefits and drawbacks of BI deployment. By understanding these limitations, we can get to expect and make more informed decisions.

8.3. Future Research Directions

Several directions appear from the empirical findings and also limitations which identified above. First, the training-validation-test gap (9 percentage points) in the query understanding module suggests that additional regularization techniques such as stochastic depth, mixout, or contrastive pre-training on enterprise domain corpora may further improve out-of-distribution generalization. Second, the expansion of the study to larger and more diverse organizational samples (including randomized or quasi-randomized assignment where feasible) would strengthen causal inference about the decision quality improvements attributable to BI deployment. Third, multimodal query understanding incorporating structured data context alongside natural language input could reduce the accuracy degradation observed for long queries by providing the model with schema-level disambiguation signals [29, 30, 31, 32, 33, 34]. Federated learning approaches for continuously updating query understanding models without centralizing sensitive enterprise query data represent a practically important direction for organizations with strict data sovereignty requirements. Finally, the integration of generative LLM capabilities going beyond query classification to full answer synthesis from structured data into the BIDQ framework would represent a natural extension of the current system toward Level 5 prescriptive analytics. The adopted methodology in this

study can be combined with techniques reported in [31, 32, 33, 34] as future trends in health, computer and telecommunications.

9. Conclusion

This paper has presented three interconnected contributions to the integration of large language models into business intelligence systems. The BI-Enabled Decision Quality (BIDQ) framework formalizes the dependency structure among data readiness, analytical capability, and decision integration, providing practitioners with a diagnostic tool for identifying the binding constraint on decision quality improvement in a given organizational context. The BERT-based query understanding module demonstrates that fine-tuning transformer encoders on enterprise query data achieves 86.5% accuracy, 86% F1-score, and 45 ms inference latency a 5.5-percentage-point accuracy improvement over the best LSTM baseline ($p < 0.001$) while remaining deployable on CPU hardware at production scale without GPU acceleration. The twelve-month field study across five industry sectors establishes that BI deployment produces average 26-percentage-point improvements in decision quality and 89% reductions in decision cycle time, with a typical financial break-even at month 12. When it comes to investing in LLM-augmented business intelligence, there are some key things to keep starters, it's crucial to figure in mind. For roadblock is out where the main - is it with the data, the analysis, or the decision-making process't identify the main constraint, you might end? If you don't up investing in the wrong area, and no matter technology is, it won't actually how advanced the improve the quality of your decisions. In other words, what does this imply for organizations? For one, they should have a lot of work to do to clear their data, particularly in the early days. It's also a good idea to focus on small, get everyone on board with the idea of investing in LLM-augmented BI. By taking a targeted approach, a sense of excitement organizations can create and anticipation for what's possible, and that can be a powerful catalyst for driving further investment and innovation. The point is to be deliberate and strategic with your investments, and only invest in the true blockages that you have. With the right approach, LLM-augmented BI can be a game-changer, but it requires careful planning, attention to detail, and a willingness the tough challenges head-on. The real-world performance of BERT-based query understanding at enterprise scale and the continuous learning pipeline for incremental retraining that saves 94% of the cost of full retraining points to the use of natural language interfaces as a primary interface technique instead of a peripheral one that augments BI systems, and therefore has the potential to significantly lower the friction of interfaces that limits decision integration. Future work should expand these results to multilingualistic and multimodal environments, validate these results in randomized organizational trials, and consider adding generative synthesis capabilities to be able to generate full answers from enterprise structured data sources instead of only classification of queries.

Acknowledgement

The author would like to thank the University of Al-Qadisiyah for providing the facilities and support that contributed to the completion of this research.

REFERENCES

1. S. Negash, and P. Gray, *Business intelligence*, In: Handbook on Decision Support Systems 2: Variations, pp. 175–193, Springer, 2008.
2. H. Chen, R. H. L. Chiang, and V. C. Storey, *Business intelligence and analytics: From big data to big impact*, MIS Quarterly, vol. 36, no. 4, pp. 1165–1188, 2012.
3. A. Gandomi, and M. Haider, *Beyond the hype: Big data concepts, methods, and analytics*, International Journal of Information Management, vol. 35, no. 2, pp. 137–144, 2015.
4. E. Brynjolfsson, L. M. Hitt, and H. H. Kim, *Strength in numbers: How does data-driven decision making affect firm performance?*, Available at SSRN 1819486, 2011.
5. J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, *BERT: Pre-training of deep bidirectional transformers for language understanding*, Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (NAACL-HLT), pp. 4171–4186, 2019.

6. A. Shollo, and R. D. Galliers, *Towards an understanding of the role of business intelligence systems in organisational knowing*, Information Systems Journal, vol. 26, no. 4, pp. 339–367, 2016.
7. W. H. Inmon, *Building the Data Warehouse*, John Wiley & Sons, 2005.
8. C. Sun, A. Shrivastava, S. Singh, and A. Gupta, *Revisiting unreasonable effectiveness of data in deep learning era*, Proceedings of the IEEE International Conference on Computer Vision (ICCV), pp. 843–852, 2017.
9. G. A. Gorry, and M. S. Scott Morton, *A framework for management information systems*, Sloan Management Review, vol. 13, no. 1, pp. 55–70, 1971.
10. R. Kimball, and M. Ross, *The Data Warehouse Toolkit: The Definitive Guide to Dimensional Modeling*, John Wiley & Sons, 2013.
11. C. Batini, C. Cappiello, C. Francalanci, and A. Maurino, *Methodologies for data quality assessment and improvement*, ACM Computing Surveys, vol. 41, no. 3, pp. 1–52, 2009.
12. R. Vidgen, S. Shaw, and D. B. Grant, *Management challenges in creating value from business analytics*, European Journal of Operational Research, vol. 261, no. 2, pp. 626–639, 2017.
13. R. Sharda, D. Delen, and E. Turban, *Analytics, Data Science, & Artificial Intelligence: Systems for Decision Support*, Pearson, 2021.
14. D. J. Power, and R. Sharda, *Decision Support Systems*, In: S. Y. Nof (Ed.), Springer Handbook of Automation, pp. 1539–1548, Springer, Berlin, Heidelberg, 2009.
15. A. Frederiksen, *Competing on analytics: The new science of winning*, Taylor & Francis, 2009.
16. T. Yu, R. Zhang, K. Yang, M. Yasunaga, D. Wang, Z. Li, J. Ma, I. Li, Q. Yao, S. Roman, et al., *Spider: A large-scale human-labeled dataset for complex and cross-domain semantic parsing and text-to-SQL task*, Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing (EMNLP), pp. 3911–3921, 2018.
17. S. Min, D. Chen, H. Hajishirzi, and L. Zettlemoyer, *A discrete hard EM approach for weakly supervised question answering*, Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP), pp. 2851–2864, 2019.
18. V. Zhong, C. Xiong, and R. Socher, *Seq2SQL: Generating structured queries from natural language using reinforcement learning*, arXiv preprint arXiv:1709.00103, 2017.
19. X. Xu, C. Liu, and D. Song, *SQLNet: Generating structured queries from natural language without reinforcement learning*, arXiv preprint arXiv:1711.04436, 2017.
20. S. F. Wamba, A. Gunasekaran, S. Akter, S. J.-F. Ren, R. Dubey, and S. J. Childe, *Big data analytics and firm performance: Effects of dynamic capabilities*, Journal of Business Research, vol. 70, pp. 356–365, 2017.
21. H. A. Simon, *A behavioral model of rational choice*, The Quarterly Journal of Economics, pp. 99–118, 1955.
22. L. Theodorakopoulos, A. Karras, A. Theodoropoulou, and C. Klavdianos, *LLMs for integrated business intelligence: A big data-driven framework integrating marketing optimization, financial performance, and audit quality*, Big Data and Cognitive Computing, vol. 10, no. 4, p. 110, 2026.
23. M. Zaharia, R. S. Xin, P. Wendell, T. Das, M. Armbrust, A. Dave, X. Meng, J. Rosen, S. Venkataraman, M. J. Franklin, et al., *Apache Spark: A unified engine for big data processing*, Communications of the ACM, vol. 59, no. 11, pp. 56–65, 2016.
24. D. Laney, *3D data management: Controlling data volume, velocity and variety*, META Group Research Note, vol. 6, no. 70, p. 1, 2001.
25. J. R. Landis, and G. G. Koch, *The measurement of observer agreement for categorical data*, Biometrics, vol. 33, no. 1, pp. 159–174, 1977.
26. I. Loshchilov, and F. Hutter, *Decoupled weight decay regularization*, arXiv preprint arXiv:1711.05101, 2017.
27. D. Bahdanau, K. Cho, and Y. Bengio, *Neural machine translation by jointly learning to align and translate*, arXiv preprint arXiv:1409.0473, 2014.
28. W. Raghupathi, and V. Raghupathi, *Big data analytics in healthcare: Promise and potential*, Health Information Science and Systems, vol. 2, no. 1, p. 3, 2014.
29. L. Wang, Y. Zhang, D. Wang, X. Tong, T. Liu, S. Zhang, J. Huang, L. Zhang, L. Chen, H. Fan, et al., *Artificial intelligence for COVID-19: A systematic review*, Frontiers in Medicine, vol. 8, p. 704256, 2021.
30. J. A. Aldhaibani, A. Yahya, R. B. Ahmad, A. S. M. Zain, M. K. Salman, and R. Edan, *On coverage analysis for LTE-A cellular networks*, International Journal of Engineering and Technology, vol. 5, no. 1, pp. 492–497, 2013.
31. J. A. Aldhaibani, A. Yahya, and R. B. Ahmad, *Optimizing power and mitigating interference in LTE-A cellular networks through optimum relay location*, Elektronika ir Elektrotechnika, vol. 20, no. 7, pp. 73–79, 2014.
32. S. Mukherjee, *Harnessing AI and Big Data for Nursing Research: Opportunities, Challenges, and Ethics*, Interdisciplinary Journal of Health, Environment and Computation, vol. 1, 2025.
33. R. Saad Mahmoud, and Q. H. Mohammed, *Investigation of Multi-Platform Media Publishing with AI Personalization Engines*, Journal of Telecommunication and Computer Technology, vol. 1, no. 1, pp. 35–48, 2026.
34. Q. H. Mohammed, S. Sadeq, M. H. Alkubaisi, and M. F. Alnaimy, *Data Mining and Machine Learning Approaches for Real Estate Valuation: A Systematic Review of Predictive Accuracy*, Journal of Telecommunication and Computer Technology, vol. 1, no. 1, pp. 57–81, 2026.