



# YÖNETİM BİLİŞİM SİSTEMLERİ ALANINDA BİLİMSEL ARAŞTIRMALAR

Editör: Dr.Öğr.Üyesi Funda AKAR

**yaz**  
yayınları

# **Yönetim Bilişim Sistemleri Alanında Bilimsel Araştırmalar**

**Editör**

Dr.Öğr.Üyesi Funda AKAR

**yaz**  
yayınları

2026

**Yönetim Bilişim Sistemleri Alanında  
Bilimsel Araştırmalar**

Editör: Dr.Öğr.Üyesi Funda AKAR

---

**© YAZ Yayınları**

Bu kitabın her türlü yayın hakkı Yaz Yayınları'na aittir, tüm hakları saklıdır. Kitabın tamamı ya da bir kısmı 5846 sayılı Kanun'un hükümlerine göre, kitabı yayınlayan firmanın önceden izni alınmaksızın elektronik, mekanik, fotokopi ya da herhangi bir kayıt sistemiyle çoğaltılamaz, yayınlanamaz, depolanamaz.

---

E\_ISBN 978-625-8574-80-7

Mart 2026 – Afyonkarahisar

Dizgi/Mizanpaj: YAZ Yayınları

Kapak Tasarım: YAZ Yayınları

YAZ Yayınları. Yayıncı Sertifika No: 73086

M.İhtisas OSB Mah. 4A Cad. No:3/3  
İscehisar/AFYONKARAHİSAR

[www.yazyayinlari.com](http://www.yazyayinlari.com)

[yazyayinlari@gmail.com](mailto:yazyayinlari@gmail.com)

## İÇİNDEKİLER

|                                                                                                                                                                                   |           |
|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----------|
| <b>Artificial Intelligence in the Public Sector as a<br/>Management Information System Issue:<br/>Perspectives on Governance, Accountability, and<br/>System Life Cycle .....</b> | <b>1</b>  |
| <i>Hakan KAYA</i>                                                                                                                                                                 |           |
| <b>Kurumsal Karar Zekâsı için Bilgi Çizgeleri: Derin<br/>Öğrenme ve GraphRAG ile Karar Destekte Yeni Nesil<br/>Mimari.....</b>                                                    | <b>29</b> |
| <i>Cevher ÖZDEN</i>                                                                                                                                                               |           |
| <b>Firmalarda Yapay Zekâ ve Büyük Dil Modellerinin<br/>Benimsenmesi: Veri Madenciliği Temelli Ülkeler<br/>Arası Bir Analiz.....</b>                                               | <b>50</b> |
| <i>Ahmet AYAZ</i>                                                                                                                                                                 |           |

*"Bu kitapta yer alan bölümlerde kullanılan kaynakların, görüşlerin, bulguların, sonuçların, tablo, şekil, resim ve her türlü içeriğin sorumluluğu yazar veya yazarlarına ait olup ulusal ve uluslararası telif haklarına konu olabilecek mali ve hukuki sorumluluk da yazarlara aittir."*

# **ARTIFICIAL INTELLIGENCE IN THE PUBLIC SECTOR AS A MANAGEMENT INFORMATION SYSTEM ISSUE: PERSPECTIVES ON GOVERNANCE, ACCOUNTABILITY, AND SYSTEM LIFE CYCLE**

**Hakan KAYA<sup>1</sup>**

## **1. INTRODUCTION**

The use of artificial intelligence in the public sector intends to enhance service delivery, streamline resource allocation, and formulate policies that are more focused on citizens (Cath et al., 2018). Nevertheless, employing AI in the public sector brings about significant issues concerning the transparency, fairness, liability and human rights aspects of the highly autonomous algorithms used for decisions in government agencies (Eubanks, 2018; Pasquale, 2015). The misuse of AI for decisions by the government can be seen in the Netherlands' 2019 SyRI (System Risk Indication) of social benefits fraud, which operated through an algorithmic model and was declared illegal by the court for discrimination and violation (van Dijck et al., 2022). In the case of the United Kingdom, an algorithm for predicting A-level results was blamed for punishing students from the economically disadvantaged schools (Ofqual, 2020). These cases demonstrate the highly risky nature of public AI systems and imply that the management of such systems should be technically, ethically, legally, and politically comprehensive. To ensure that public AI systems are trustworthy, just, and

---

<sup>1</sup> Dr., T.C. Kestel Municipality, ORCID: 0000-0002-0812-4839.

democratically verifiable, AI governance and algorithmic accountability are paramount. According to this reassessment, the main shortcomings that impede responsible AI are non-technical challenges, however, they have been addressed in such a way as to be applicable in this new algorithmic environment. These problems include: (i) lack of consideration for user-centered design in the development lifecycle; (ii) ambiguity in IT governance and liability (the "accountability gap"); (iii) institutional competency; and (iv) absence of community participation in system design.

To solve these fundamental issues, the paper ends with a number of policy and operational suggestions that are informed by MIS. Some of these include setting up organizational units for the independent auditing of algorithms; carrying out participatory design and agile development methods in public AI projects; encouraging open source public software as a way of raising its transparency; and establishing sandboxes for controlled experiments.

Ultimately, however, this chapter asserts that the success of AI within the public sector shall be assessed not by its mathematical accuracy but by the quality of human and institutional systems that govern it—and by that token, for academics and developers within the realms of public engineering and data ethics, a crucial concern for public administration and strategic management of information systems within the 21st century shall be made. This chapter offers both theoretical and practical contributions to the literature by reframing public AI systems with MIS concepts such as system development life cycle (SDLC), IT governance, and user-centered design, and also a conceptual and analytical framework that can be operationalized and tested by future empirical research.

## **2. BASIC CONCEPTS**

### **2.1. Artificial Intelligence Governance (AI Governance)**

AI governance is the entirety of rules, principles, policies, and institutional mechanisms that encompass the design, development, deployment, and monitoring processes of AI systems (Cath, 2018). In the public sector, AI governance is based on three fundamental principles:

- **Transparency:** Understandability of how the systems work.
- **Justice:** Prevention of discrimination between individuals or groups.
- **Auditability:** Decisions must be reversible and subject to appeal.

In contrast to other existing studies of governance of AI, this chapter reframes public AI systems as artifacts of MIS and examines failures based on misalignment of system type and governance.

### **2.2. Algorithmic Accountability**

Citron and Pasquale (2014) define algorithmic accountability as a framework that ensures “algorithms are traceable, auditable, and accountable in accordance with legal, ethical, and social norms.” This concept includes the following components:

- **Data responsibility:** Representation of training data, bias analysis, and updating.
- **Model explainability:** The ability to express decision logic technically and verbally.

- Appeal mechanisms: The ability of citizens to appeal algorithmic decisions.
- Audit and evaluation: Periodic algorithmic impact assessment by independent bodies.

The technical success of AI systems is directly related not only to algorithmic accuracy but also to the adoption of these systems by public employees (User Acceptance). From the perspective of the Technology Acceptance Model (TAM) in the YBS literature, mistrust of ‘algorithmic decision-makers’ in public bureaucracy or fear of job loss can lead to system failure. Therefore, to guarantee democratic legitimacy of Algorithmic Impact Assessment (AIA) processes, it is necessary to have not just technical auditors but also end users like public officials and citizen representatives.

### **2.3. High-Risk Artificial Intelligence Systems**

Apart from that the European Union Artificial Intelligence Act or AI Act states that AI systems used by the public sector are labeled as "high-risk" (European Commission, 2024).

An example of systems that would be regarded as this category are:

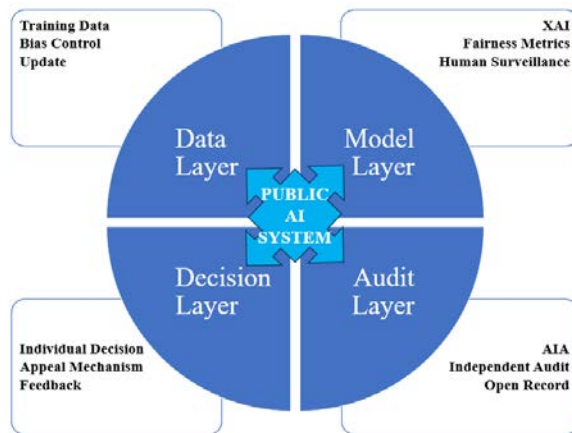
- Automatic assessment of social assistance, tax, or immigration applications,
- Policing and judicial prediction systems,
- Healthcare access allocation.

These systems are subject to mandatory compliance assessment, quality data management, human oversight, and session recording requirements.

## **2.4. AI Governance in the Public Sector from a Management Information Systems (MIS) Perspective**

The discipline of Management Information Systems focuses on managing information technologies in alignment with organizational goals and processes. AI systems in the public sector are an evolved and more complex form of traditional MIS. In this context, AI governance and accountability intersect directly with topics that are core areas of interest for MIS, such as system lifecycle management, IT governance frameworks (e.g., COBIT), data management, decision support system design, and the assessment of the organizational/social impacts of information systems.

The multi-layered framework presented in Figure 1 demonstrates that an AI system should be treated as an end-to-end MIS project; this project encompasses not only the software development process but also classic MIS processes such as data infrastructure, user acceptance, continuous monitoring, evaluation, and corporate risk management.



**Figure 1. Multi-layered algorithmic accountability framework in public AI systems**

The Data Layer is the base of the system and what it seeks to achieve in essence is the quality, neutrality and the timeliness of the information the AI model will be learning from. This layer requires disclosing the sources of the training data in a very open manner, keeping an eye on biases that might be present in the data, and making sure there are regular data refresh cycles to keep up with the changing world scenarios. It goes without saying that a strong data layer is needed for the overall system to be fair and accurate.

The Model Layer is designed to help the user understand how the algorithm works and it is also charged with performance monitoring at all times. Explainable Artificial Intelligence (XAI) methods increase users' confidence in the model by revealing the thought process behind the predictions made. Checking the model with regular metrics for performance as well as fairness standards makes it possible to recognize situations of unjust discrimination or privilege. In addition, the use of human-in-the-loop techniques is like having a safety net, it compensates for the deficiencies of the automation by giving the expert human the possibility and power to intervene and approve decisions at the most critical points.

The Decision Layer deals with the ways individuals receive the model's outputs and what rights these individuals have. In this layer, communicating individual decision outputs to a single person in clear and simple language is a must. Most importantly, it is the right of the individual to allow him or her access to a transparent appeal mechanism when the decision is unfavorable.

Such appeals and the user feedback that is obtained constitute a very valuable resource for system enhancement, thus unleashing a continuous feedback loop that brings information to both the Data and Model layers.

The Audit Layer is intended to be at the system's regulatory and ethical level. The primary aim of the AIA tool has been to analyze even before the release, the potential risks as well as the social impacts of the system. At every step, impartial independent auditors' reviews help increase transparency as well as accountability. The open algorithmic record, which is a result of these documentation processes and the system's main features, acts as a public record of the system operations hence the highest level of accountability.

The multilayered framework illustrated in Figure 1 enables us to understand that a fair and responsible artificial intelligence system is not just a technical model. Rather, it is a complex and dynamic living ecosystem consisting of processes that continuously interact with each other, starting from data and going all the way to social oversight.

### **3. THEORETICAL AND POLICY FRAMEWORKS**

#### **3.1. International Policy Initiatives**

##### **3.1.1. European Union Artificial Intelligence Act (AI Act)**

The AI Act was officially confirmed by the European Parliament and Council in the final form on March 13, 2024 (Regulation (EU) 2024/1689), but the introduction of the rules for high-risk systems will only be gradually from 2026 (European Commission, 2024).

The high-risk category is defined in Annex III and contains social assistance, migration, law enforcement prediction, and healthcare access in public administration; it makes conformity assessment, data governance, human oversight, and technical documentation requirements mandatory.

Figure 2 shows a hierarchy or pyramid model that classifies algorithmic systems by their potential threat levels and is consistent with the European Union's Artificial Intelligence Act (AI Act). This model has four main levels, where risk and the corresponding regulatory intervention increase step by step.



**Figure 2. EU AI Act risk classification and the place of public AI systems**

Minimal Risk is at the base of the pyramid, covering the greatest area and the lowest level.

For example, spam filters are systems which carry an acceptable level of risk and are generally assessed under the current legislation. Systems of Limited Risk, which is one level above (for example, chatbots), have clear transparency requirements so that users have the right to know that they are interacting with an artificial intelligence.

Most of the attention and detail of the model goes to the High Risk category. These are systems that can significantly impact, either positively or negatively, the fundamental rights or safety of persons, and their use in the public sector is especially significant. It identifies four main areas of social risk where this is most evident: distributing social benefits, deciding on immigration, predicting crime (policing prediction), and providing healthcare. Decisions in these sectors have an immediate influence on people and in some cases these effects may be irreversible. Therefore, these systems are subjected to

stricter regulation than others. Furthermore, a note at the bottom of the picture lists the main rights protection measures for such systems: use of high-quality, unbiased data; guarantee of effective human oversight (human-in-the-loop); detailed technical documentation; and suitability test.

Last, the very top of the pyramid is reserved for those practices that are against democratic societal values and carry risks that are unacceptable. The proposal of such systems (e.g., state-imposed social scoring mechanisms) is prohibited by law.

Figure 2 illustrates the need for a risk-based, proportionate, and multi-layered regulatory framework for various algorithmic systems rather than a "one size fits all" approach.

It clearly conveys that as risk increases, not only transparency, accountability, and human oversight mechanisms need to be complemented by more and more rigorous measures but also that these elements must be improved first and foremost.

### **3.1.2. Other International Developments**

The Canadian Government in 2019 mandated an AIA policy to all federal agencies, requiring them to disclose and evaluate their algorithmic decision systems. The evaluation documentation reveals various aspects of the system's data sources, the consideration of biases, error rates, and the procedure by which citizens may challenge the decision. For instance, the Ministry of Employment and Social Development's system for prioritizing unemployment benefits was retargeted at the time of the AIA process as it was found to be historically biased against female applicants. The goal of the method is to make technical and ethical oversight permanent features of the institution (Government of Canada, 2023).

The UK Information Commissioner's Office (ICO) together with the Alan Turing Institute ran a pilot program of the "Algorithmic Transparency Sandbox" for public institutions in 2022, 2023. This framework permitted local governments to deploy experimental high-risk algorithmic systems (e.g., social assistance allocation or housing allocation models) without risk to the public. The participating organizations used XAI techniques such as SHAP and gathered the citizens' feedback on the clarity of these explanations. Consequently, some of the participants who took part in the pilot decided to stop using indirect discrimination indicators such as "postal codes" (ICO & The Alan Turing Institute, 2023).

In 2021, the Netherlands initiated the Amsterdam and Helsinki Algorithm Register to promote public algorithm transparency. The public is allowed to see how certain algorithms are used to make decisions, what data was relied on, and which institution is responsible. Behind this initiative lies the SyRI (System Risk Indication) system, which was declared illegal by the court in 2019. SyRI targeted low-income neighborhoods to detect social assistance fraud and was overturned by the Constitutional Court on grounds of "violation and discrimination." This incident demonstrated why algorithmic accountability is a corporate imperative (van Dijck et al., 2022; Algorithm Register, 2023).

France has mandated the disclosure of public algorithms through the Etalab platform to increase transparency in digital services. Since 2021, algorithmic decision systems used by state institutions have been registered with Etalab, along with their purpose, data source, performance metrics, and responsible person information. This registration system is being brought into line with the EU AI Act (Etalab, 2023).

In line with global developments, Turkey has established a roadmap for the use of AI in public administration with its 2021-2025 National Artificial Intelligence Strategy. However, considering the strict regulations introduced by the EU AI Act, the compliance of public projects in Turkey with the Personal Data Protection Law (KVKK) constitutes a critical threshold in terms of algorithmic accountability (Official Gazette, 2016). In particular, how the balance between ‘explicit consent’ and ‘legitimate interest’ in public institutions’ data processing activities can be protected with non-transparent algorithms is an issue that needs to be discussed further in the local literature.

Table 1 presents a comparative summary of the different AIA approaches applied to algorithmic systems in the public sector across countries and regions. It is based on four main criteria which comprise whether there is a legal requirement for an AIA regime, the scope of the regime, the use of open logging as a transparency measure, and the involvement of independent oversight.

**Table 1. Comparison of International AIA Applications**

| Country / Region | Mandatory ? | Scope                            | Open Registration ?           | Independent Audit | References                             |
|------------------|-------------|----------------------------------|-------------------------------|-------------------|----------------------------------------|
| Canada           | Yes         | Federal algorithmic systems      | No                            | Partially         | Government of Canada (2023)            |
| England          | No          | Volunteer sandbox participants   | Yes                           | Yes               | ICO & The Alan Turing Institute (2023) |
| EU (AI Act)      | Yes         | High-risk systems                | Recommended                   | Compulsory        | European Parliament (2024)             |
| Netherlands      | Partially   | Municipal level                  | Yes                           | Yes               | Algorithm Register (2023)              |
| France           | Yes         | National digital services        | Yes                           | It is planned     | Etalab (2023)                          |
| Türkiye          | No          | No (there are local initiatives) | Partially (data transparency) | No                | Municipalities' open data portals      |

### **3.2. Theoretical Approaches**

#### **3.2.1. Value-Sensitive Design**

Friedman and Kahn (2003) suggest that ethical values should be incorporated into the technology design process. The values of fairness, privacy, and participation should be considered at the coding stage itself if public AI systems are involved.

#### **3.2.2. Digital Public Service Ethical Framework**

Sargeant (2025) reveals that public AI systems need to be compatible with the concept of "public good." This model entails measuring the degree to which algorithms exacerbate societal inequalities instead of simply fixing inefficiencies.

## **4. IMPLEMENTATION MECHANISMS AND REAL-WORLD EXAMPLES**

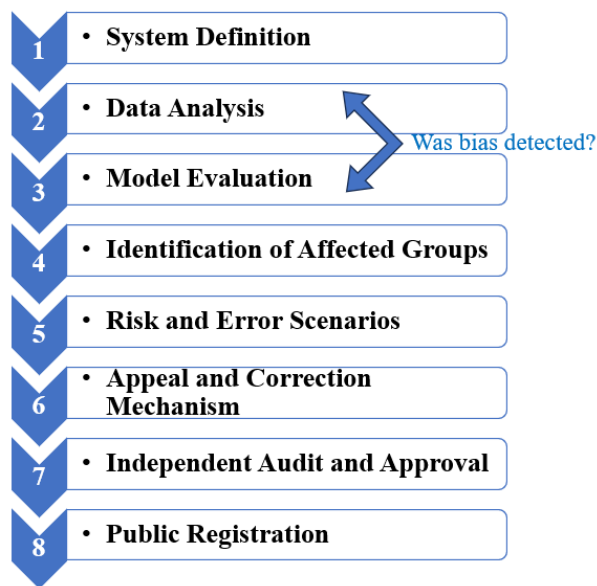
The implementation mechanisms discussed in this section are, in fact, the concrete counterparts of effective MIS design and management in the context of public AI. Algorithmic Impact Assessment can be seen as a system analysis and design tool, open records as an information system documentation and transparency mechanism, and XAI as a decision support system user interface and reporting component.

### **4.1. Algorithmic Impact Assessment (AIA)**

Canada has mandated AIA for all federal algorithmic systems since 2019 (Government of Canada, 2023). The form includes the following questions:

- What data does the system use?
- Which groups could be negatively impacted?
- What is the error rate?
- Is there an appeal mechanism?

Iterative AIA processes for high-risk systems (Figure 3), XAI integration (LIME, SHAP), and transparent algorithmic records are critical tools. While Amsterdam and Helsinki provide pioneering examples of documentation, the revelation in 2025 that Amsterdam's social assistance algorithm still contained discrimination and the subsequent suspension of the pilot demonstrate that transparency alone does not guarantee fairness (Guo et al., 2025). The process includes eight main steps that focus first on a complete technical review of the system and then also on its social ethics and administrative accountability.



**Figure 3. Standardized Algorithmic Impact Assessment (AIA) flow for public AI systems**

The method is a series of eight steps that thoroughly evaluate a system comprehensively beyond just the technology aspects, also including the social ethics and the administrative accountability dimensions.

The initial step is System Definition and it is all about clarifying and documenting the use, the extent, and the kind of decisions the system will be making entirely and very clearly. Then, the following step, Data Analysis, is focused on doing very detailed investigations such as checking the original source of the data sets to be used, their representativeness of the society, and the presence of possible biases.

Model Evaluation gauges both the technical performance of the system (accuracy) and the system's fairness (using metrics such as demographic parity and equal opportunity). The figure particularly stresses that a problem revealed at this stage (e.g., a “Yes” answer to the question “Was bias detected?”) will lead to a corrective action by redirecting the process to the Data Analysis or model retraining stage through a cyclical feedback loop.

The phase of the process that concentrates on the system's social impacts after the technical evaluation is done. In the Identification of Affected Groups phase, potential social, economic, or gender-based unequal impacts of the decision are identified. The Risk and Error Scenarios section discusses the real-life consequences and injustices that even technical errors (false positives/negatives) can cause.

Therefore, the design of the Appeal and Correction Mechanism (citizen appeal route) is based on these evaluations.

The last three steps are targeted at confirming the system's final approval and public accountability. The independent audit and approval step means that the ethics committee or an accredited external auditor prepare an independent report. The system that has successfully gone through all these stages is finally recorded in the Public Registry as a sign of transparency and is included in the open algorithm ledger, thus, being available for public scrutiny.

Consequently, Figure 3 illustrates a non-linear checklist but a vibrant and introspective process. Its "Yes/No" decision points with circular feedback arrows clearly indicate that this audit tool is essentially a self-correcting quality assurance system that must return the product for revision and improvement when issues arise or are detected.

Usually, the 'Deployment' phase is seen as the climax of the whole cycle in the traditional SDLC models. However, for AI projects, this phase unveils a new beginning due to 'Continuous Learning' and 'Data Drift'. Introducing 'Continuous Learning' and 'Data Drift' in AI projects take the 'Deployment' phase in traditional SDLC to a new beginning rather than the end of the cycle. Community use of AI raises the risk of algorithmic bias that drifts over time and consequently necessitates a 'continuous audit' mechanism at every stage of the SDLC system. This is a testament that an information system should even be viewed not just as a technical tool but also as a living managerial process.

#### **4.2. Algorithmic Registries**

In 2021, the Municipality of Amsterdam organized workshops involving not only data scientists but also social service workers, applicants, and NGOs to develop an algorithm for prioritizing social assistance applications. During this process, it was decided to integrate multidimensional indicators such as “number of children + housing conditions” into the system instead of “income thresholds.” As a result, the algorithm gained not only technical accuracy but also social legitimacy (Algorithm Register, 2023). Citizens can see which systems are used in which decisions, the data sources, and the responsible institution. This is an application that supports participatory oversight.

### **4.3. Explainable Artificial Intelligence (XAI) Integration**

XAI techniques in public AI systems increase the comprehensibility of decisions:

- LIME (Local Interpretable Model-agnostic Explanations): Explains feature importance for individual decisions (Ribeiro et al., 2016).
- SHAP (SHapley Additive exPlanations): Provides consistent contribution analysis based on game theory (Lundberg & Lee, 2017).

In 2020, the Danish Social Security Agency developed an open-source module to provide SHAP explanations to citizens regarding social assistance applications. The code was published on GitHub and adapted by Sweden and Norway. This allowed the three countries to share the same fairness metrics (De Bruijn et al., 2022).

### **4.4. Management Information Systems (MIS)-Based Public AI Application Scenarios and Practical Tips**

The earlier parts of the text were about the theory and legal aspects. Here, in this part, we talk about the examples and useful pieces of advice relating to the basic principles and techniques of the MIS field that can be realized in the creation and operation of public AI systems. The idea is that public AI projects shall not be only technology installations but also the use of strategic information system initiatives that are in line with the goals of the organization.

#### **4.4.1. Application of Participatory Design and Agile Methodologies**

MIS Principle: System Development Life Cycle (SDLC), User-Centered Design, Agile Project Management.

Scenario: A municipality wants to develop an algorithmic decision support system to be used in social housing allocation. Instead of the traditional “waterfall” model, a participatory and agile approach is adopted.

#### Practical Implementation Steps and MIS Links:

1. Stakeholder Mapping and Needs Analysis (MIS: System Analysis): A stakeholder group is formed, consisting not only of technical staff and managers, but also social service specialists, civil society organization (CSO) representatives, citizens who have previously applied for housing, and lawyers. This group participates in structured workshops during which the system's requirements are established (Zuiderwijk et al., 2021).

2. Agile Sprints and Prototyping (MIS: System Design and Prototype): Development is planned in two-week sprints. After each sprint, a working prototype (e.g., a simple priority scoring interface) is demonstrated to the stakeholder group. The next sprint backlog consists of feedback.

3. Usability and Acceptance Testing (MIS: System Testing and Acceptance): The final system is tested for usability with the involvement of the target users (social service workers). Besides, a moral and legal acceptance test is performed to ensure the fairness and transparency of the system's outputs.

Output: A system that is socially acceptable, user-friendly, and increases public trust besides being technically correct.

#### **4.4.2. Organizational Structure for Independent Algorithmic Audit**

MIS Principle: IT Governance (COBIT/ITIL), Information Systems Audit, Three-Line Defense Model.

Scenario: A country intends to create an independent review body for public AI systems with a high level of risk.

### Practical Implementation Steps and MIS Links:

1. Establishment of the Algorithmic Audit Authority (MIS: IT Governance Structure): A "National Algorithmic Audit Authority" would be created either through incorporation into existing bodies, eg the Court of Auditors or the Data Protection Authority, or as a completely independent entity. Such an authority would host a multidisciplinary team of auditors, data scientists, lawyers, and ethics specialists (De Vries et al., 2016).

2. Aligning the Audit Framework with MIS Frameworks: The audit process is adapted to COBIT's "Command, Execute, Control, Evaluate" cycle or ITIL's service lifecycle. For example, the audit scope includes:

- Governance: Project management documentation, RACI matrix, risk assessment reports.
- Data and Model: Independent verification of the AIA report, bias tests on data sets, evaluation of model performance and fairness metrics.
- Operations and Support: Functionality of appeal mechanisms, session logs, testing of human oversight protocols.

3. Ex-Ante and Ex-Post Audit Cycle: The system undergoes mandatory pre-deployment (ex-ante) auditing. Post-deployment, performance and ethical compliance audits are repeated at regular intervals (ex-post).

Output: A clear chain of responsibility, increased accountability, and an impartial technical intermediary that builds trust between citizens and institutions.

#### **4.4.3. Open Source Public Software and Transparent Infrastructure Development**

MIS Principle: Information Systems Architecture, Procurement Strategy, System Documentation.

Scenario: The Ministry of Health has developed a model that predicts hospital emergency room workload. It wants to increase the transparency and reusability of this model.

Practical Implementation Steps and MIS Links:

1. Establishing an Open Source Policy (MIS: IT Strategy and Policy): A guideline is published requiring the source code of publicly funded AI software to be released under open licenses (e.g., MIT, GPL), except for legitimate exceptions such as trade secrets or national security (Janssen & Kuk, 2021).

2. Modular and Documented Architectural Design (MIS: System Architecture): The system is developed in independent modules such as data cleaning, feature extraction, model training, and API layer. Technical and functional documentation is prepared for each module.

3. Open Source Platform and Community Management: The code is published in a government-owned Git repository or on GitHub. A “Contribution Guide” is created to encourage contributions from other public institutions, universities, and volunteer developers. A management process is defined for bug reports and improvement suggestions.

4. Regulation of Tender and Procurement Criteria (MIS: Procurement Management): Criteria encouraging bidding companies to offer open source solutions or evaluating contributions to open source are added to public tender documents.

Output: Lower maintenance costs, higher security.

#### **4.4.4. Controlled Pilot Application Management with Algorithmic Sandboxes**

MIS Principle: Project Risk Management, Pilot Application, User Acceptance Testing (UAT).

Scenario: The Ministry of Interior plans to implement a text analysis model to be used in the preliminary review of immigration applications.

Practical Implementation Steps and MIS Links:

1. Sandbox Environment Setup (MIS: Test Environment Management): A secure digital sandbox environment is created that is isolated from the production environment but simulates the actual data structure and flow. All data is anonymized and sensitive information is removed.

2. Controlled Pilot and Feedback Loop (MIS: User Acceptance Testing): The model is tested in the sandbox for a limited period (e.g., 3 months) and with a limited user group (e.g., 20 trained officers and 50 volunteer applicants). It runs on a copy of real applications, but the outputs do not affect actual decisions.

3. Risk and Impact Monitoring (MIS: Project Risk Monitoring): During the sandbox period, qualitative metrics such as technical errors (false positives/negatives), user trust in the system, comprehensibility of decisions, and satisfaction are monitored. Unanticipated social effects are quickly noticed (De Bruijn et al., 2022).

4. “Go/No-Go” Decision Point: After the experiment, a thorough evaluation is done based on a report on the technical performance, an ethical assessment, and feedback from the users. This evaluation leads to a decision by the project management as to whether the project should be carried out at a larger scale (Go), modified through changes (Iterate), or discontinued (No-Go).

Output: The risk of a big failure resulting in a loss of trust of the public is lowered as now a more advanced system has been through testing with real data.

Even though theoretical frameworks and regulatory models potentially laying out the governance requirements for public AI systems have been described, they are mostly not fully practical and thus complete 'roadmaps' of how to govern such systems. In order to bridge this gap, as well as to help the operational counterparts of the principles derived from the MIS discipline become more tangible, the table below arranges the mapping of the four fundamental MIS processes to the critical implementation tools in public AI governance in a matrix format. Table 2 is basically a strategic 'inform and guide' tool for decision-makers and practitioners drawn from a summarization of concrete mechanisms applied to the implementation of each principle alongside their expected major advantages, and the present academic studies that back these methods.

**Table 2. Summary Matrix of MIS-Based Public AI Application Scenarios**

| MIS Principle / Process           | Application Tool in the Context of Public AI | Key Benefits                                                                                   | References               |
|-----------------------------------|----------------------------------------------|------------------------------------------------------------------------------------------------|--------------------------|
| Agile SDLC & Participatory Design | User Workshops, Prototype Testing            | Increased user acceptance, fewer requirement errors, social legitimacy                         | Zuiderwijk et al. (2021) |
| IT Governance & Audit             | Independent Algorithmic Oversight Authority  | Net responsibility, compliance assurance, alignment with the three-line defense model          | De Vries et al. (2016)   |
| System Architecture & Procurement | Open Source Public Software Policy           | Transparency, auditability, reusability, cost efficiency, ecosystem development                | Janssen & Kuk (2021)     |
| Project Risk Management & UAT     | Algorithmic Sandbox                          | Early risk detection, controlled testing, realistic user feedback, “go/no-go” decision support | De Bruijn et al. (2022)  |
| Decision Support Systems          | Automatic AIA Reporting Dashboards           | Real-time impact monitoring, data-driven decision-making, streamlining governance processes    | Reisman et al. (2018)    |
| Human Resource Planning           | Public AI Competency Matrix and Academy      | Corporate capacity development, closing the skills gap, sustainable expertise management       | Amirova et al. (2025)    |

Table 3 illustrates a categorization which shows the conversion of conventional YBS structures into ‘intelligent

systems' by means of AI. For instance, a traditional Decision Support System (DSS) normally depends on fixed rules, whereas an AI-enabled DSS changes its framework dynamically through the use of predictive models (ML) that get trained on historical data. At this point, the MIS discipline should focus not only on the system's output but also on the controllability of the 'decision-making logic' (XAI - Explainable AI). This is an operational necessity to overcome the 'black-box' problem in public administration.

**Table 3. MIS-Based Classification of Public Artificial Intelligence Applications, Risks, and Governance Tools**

| MIS Type                                                             | Public Use Area                                     | Basic Risk Profile                                          | Appropriate Governance Tool                                              |
|----------------------------------------------------------------------|-----------------------------------------------------|-------------------------------------------------------------|--------------------------------------------------------------------------|
| Decision Support Systems (DSS) / Executive Information Systems (EIS) | Social assistance and welfare policies              | Discrimination, data representation issues, automation bias | Algorithmic Impact Assessment (AIA), fairness metrics, human-in-the-loop |
| Transaction Processing Systems (TPS) + DSS (Hibrit)                  | Immigration, asylum, and visa assessment            | Uncertainty of liability, difficulty of legal challenge     | Open algorithm registry, role-based IT governance, appeal mechanisms     |
| Strategic Information Systems (SIS)                                  | Proactive policing, security, and crime analytics   | Feedback loops, structural bias, loss of legitimacy         | Strategic audit, performance indicators, public transparency reports     |
| Clinical Decision Support Systems (CDSS)                             | Access to and prioritization of healthcare services | Suppression of clinical intuition, erosion of trust         | Explainable AI (XAI), user-centered interface design, ethics committees  |

This section contributes to the literature by addressing public AI not as individual algorithms, but as socio-technical systems that are deeply intertwined in the different types of MIS. Looking into the algorithmic risks that are specific to the system type is one of the ways that humanizes the efforts of AI governance with the MIS lifecycle management. The research has come up with a highly generative theoretical framework that can diagnose the reasons of breakdown of public AI projects on the

basis of MIS displacement rather than model accuracy, and also offers a classification which can be tested in future empirical research.

## **5. CONCLUSION AND EVALUATION**

The chapter argues that the biggest barriers to public sector AI are not tied to the mathematical aspects or the intricacies of AI models. Instead, they stem from AI systems being treated as technology systems, separate from the fact that they are high-stakes information systems. If AI systems are viewed in isolation, the typical failure manifestations in terms of transparency, accountability, legitimacy, and performance can be seen. The findings indicate that public sector AI governance should not be seen as an additional issue related to ethics or a mere compliance with regulations, but rather as a fundamental element of MIS.

On the conceptual front, the paper builds on prior works and shows that lack of transparency, discrimination, and irresponsibility lead to failure of the system, low user acceptance, and misallocation of control, respectively. The continued relevance of concerns over “black boxes” reflects more than the need to conceal the workings of a particular model; rather, it indicates more profound issues concerning disparities between the outputs of system decisions and decision contexts and requirements. Essentially, this represents issues of information quality equally as much as anything does in accordance with classical MIS theories.

Similarly, the concern about the accountability gap indicates that the classical limit of IT governance applies to semi-autonomous decision systems as well. In these circumstances, the distribution of responsibility among developers, data providers, and government representatives leads to a very common MIS governance problem: confusion about roles within a broader

decision-making process. Therefore, in this chapter, accountability of an algorithm as a phenomenon not only at decision but through multi-levels of design, implementation, observation, and post-decision stages is presented.

The dialogue also insists that most of the issues related to AI in the public sector arise due to institutional capacity constraints rather than being technically unfeasible. A lack of human capital, data governance, and algorithm literacy among decision-makers is often responsible for inefficient management and might often lead to overdependence on AI results. The answer to this challenge is more than just incorporating more technology into public sector MIS plans.

At the level of management and policy, one could deduce that deployment of AI in public setting is successful when alignment of governance along with positioning have been done properly. The type of system, the level of risk exposure, and the criticality of a decision would determine how the governance structures are changed. This can be seen as an extension of contingency theory in algorithmic management.

Participatory processes are missing in old-school SDLCs where optimization is more focused on technical velocity rather than how people and organizations work. If we consider public sector AI as an MIS artifact, then it is worthwhile to acknowledge the role of participatory design in this regard.

Keeping the fusion of theory and practice in mind, this wrap-up discussion places AI governance as the natural, though changing, extension of the MIS studies field. Henceforth, empirical studies can adopt this framework to explore how different MIS structures may affect the societal and organizational benefits from AI in the public sector.

## REFERENCES

- Algorithm Register. (2023). *Amsterdam/Helsinki algorithm register*. <https://algorithmeregister.amsterdam.nl/>
- Amirova, A., Iskenderova, S., Zhanseitov, A., Daueshova, A., & Zhunissova, Z. (2025). Does Government AI Readiness Improve Human Resource Management Efficiency in The Public Sector?. *Journal of Cultural Analysis and Social Change*, 1160-1176. <https://doi.org/10.64753/jcasc.v10i4.2995>
- Cath, C. (2018). Governing artificial intelligence: Ethical, legal, and technical opportunities and challenges. *Philosophical Transactions of the Royal Society A*, 376(2133), 20180080. <https://doi.org/10.1098/rsta.2018.0080>
- Cath, C., Wachter, S., Mittelstadt, B., Taddeo, M., & Floridi, L. (2018). Artificial intelligence and the ‘good society’: The US, EU, and UK approach. *Science and Engineering Ethics*, 24, 505–521. <https://doi.org/10.1007/s11948-017-9901-7>
- Citron, D. K., & Pasquale, F. (2014). The scored society: Due process for automated predictions. *Washington Law Review*, 89(1), 1–33.
- De Bruijn, H., Warnier, M., & Janssen, M. (2022). The perils and pitfalls of explainable AI: Strategies for explaining algorithmic decision-making. *Government Information Quarterly*, 39(1), 101666. <https://doi.org/10.1016/j.giq.2021.101666>
- De Vries, H., Bekkers, V., & Tummers, L. (2016). Innovation in the public sector: A systematic review and future research agenda. *Public administration*, 94(1), 146-166. <https://doi.org/10.1111/padm.12209>

- Etalab. (2023). *Référentiel des algorithmes de l'État*. <https://www.etalab.gouv.fr/referentiel-des-algorithmes-de-letat/>
- Eubanks, V. (2018). *Automating inequality: How high-tech tools profile, police, and punish the poor*. St. Martin's Press.
- European Commission. (2024). *Proposal for a regulation laying down harmonised rules on artificial intelligence (Artificial Intelligence Act): Final agreed text (March 2024)*. <https://digital-strategy.ec.europa.eu/en/library/proposal-regulation-laying-down-harmonised-rules-artificial-intelligence>
- European Parliament & Council of the European Union. (2024). *Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 March 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act)*. Official Journal of the European Union, L1, 1–144. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>
- European Parliament. (2024). *Regulation (EU) 2024/1689 of the European Parliament and of the Council on laying down harmonised rules on artificial intelligence (Artificial Intelligence Act)*. <https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>
- Friedman, B., & Kahn, P. H. (2003). Human values, ethics, and design. In J. Jacko & A. Sears (Eds.), *The human-computer interaction handbook* (pp. 1241–1266). Lawrence Erlbaum.
- Government of Canada. (2023). *Directive on Automated Decision-Making (Version 2.0)*. Treasury Board of Canada Secretariat. <https://www.tbs-sct.canada.ca/pol/doc-eng.aspx?id=32592>

- Guo, E., Gabriel, G., & Braun, J. C. (2025). Inside Amsterdam's high-stakes experiment to create fair welfare AI. *MIT Technology Review*.
- Information Commissioner's Office (ICO) & The Alan Turing Institute. (2023). *Explaining algorithmic decisions: Final report from the AI and data protection sandbox*. <https://ico.org.uk/media2/migrated/4024793/good-with-sandbox-exit-report.pdf>
- Janssen, M., & Kuk, G. (2016). The challenges and limits of big data algorithms in technocratic governance. *Government Information Quarterly*, 33(3), 371-377. <https://doi.org/10.1016/j.giq.2016.08.011>
- Lundberg, S. M., & Lee, S. I. (2017). A unified approach to interpreting model predictions. *Advances in Neural Information Processing Systems*, 30, 4765–4774.
- Official Gazette. (2016, 24 Mart). 6698 sayılı Kişisel Verilerin Korunması Kanunu uygulama yönetmeliğinde değişiklik. Sayı: 29677. <https://www.resmigazete.gov.tr/eskiler/2016/04/20160407-8.pdf>
- Ofqual. (2020). *Review of the summer 2020 exam grading process*. Office of Qualifications and Examinations Regulation. <https://www.gov.uk/government/publications>
- Pasquale, F. (2015). *The black box society: The secret algorithms that control money and information*. Harvard University Press.
- Reisman, D., Schultz, J., Crawford, K., & Whittaker, M. (2018). Algorithmic impact assessments: a practical Framework for Public Agency. *AI Now*, 9.
- Ribeiro, M. T., Singh, S., & Guestrin, C. (2016). "Why should I trust you?": Explaining the predictions of any classifier.

*Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 1135–1144. <https://doi.org/10.1145/2939672.2939778>

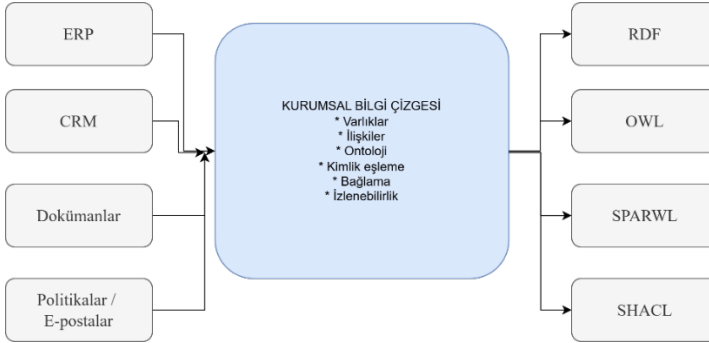
- Sargeant, H. (2025). From Estimation to Discrimination: Algorithmic Bias, Predictive Uncertainty, and Anti-Discrimination Law. *University of Cambridge Faculty of Law Research Paper*, (11). <http://dx.doi.org/10.2139/ssrn.5283387>
- van Dijck, J., Poell, T., & de Waal, M. (2022). *The platform society as a contested domain: How public values can be reclaimed*. Oxford University Press.
- Zuiderwijk, A., Chen, Y. C., & Salem, F. (2021). Implications of the use of artificial intelligence in public governance: A systematic literature review and a research agenda. *Government information quarterly*, 38(3), 101577. <https://doi.org/10.1016/j.giq.2021.101577>

# KURUMSAL KARAR ZEKÂSİ İÇİN BİLGİ ÇİZGELERİ: DERİN ÖĞRENME VE GRAPHRAG İLE KARAR DESTEKTE YENİ NESİL MİMARİ

Cevher ÖZDEN<sup>1</sup>

## 1. GİRİŞ

Kurumsal Karar Zekâsı (Enterprise Decision Intelligence), klasik Karar Destek Sistemleri'nin yapay zekâ, ileri analitik, otomasyon ve sürekli geri besleme mekanizmalarıyla genişleyen çağdaş bir evresini ifade etmektedir. Bu yaklaşımda kararlar, yalnızca belirli bir anda verilen çıktılar olarak değil, tasarlanabilen, izlenebilen ve kurumsal performans açısından yönetilebilen stratejik varlıklar olarak ele alınmaktadır. Bu bölümde, kurumsal karar kalitesini artıran temel bileşen olarak bilgi çizgeleri (knowledge graphs) “anlamsal omurga” işleviyle konumlandırılmaktadır (Şekil 2).



**Şekil 2. Kurumsal bilgi çizgesi: Veri bütünleştirme**

<sup>1</sup> Dr. Öğr. Üyesi, Çukurova Üniversitesi, Fen-Edebiyat Fakültesi, Bilgisayar Bilimleri Bölümü, ORCID: 0000-0002-8445-4629.

Kurumsal bilgi çizgeleri, ERP ve CRM gibi yapılandırılmış veri kaynakları ile dokümanlar, e-postalar ve politika metinleri gibi yapılandırılmamış içerikler arasında kimlik, bağlam ve ilişki sürekliliği kurarak veri silolarının azaltılmasına katkı sunmaktadır (Kostopoulos ve ark., 2026). Bu çerçevede RDF ve OWL gibi standartlar aracılığıyla anlamsal modelleme yapılmakta, SHACL ile veri kalitesi ve doğrulama sağlanmakta, SPARQL ile ise sorgulanabilir ve güvenilir bir kurumsal bilgi katmanı oluşturulmaktadır. Söz konusu katman üzerinde çalışan GNN tabanlı çizge öğrenme yaklaşımları, bağlantı kestirimi (link prediction) ve düğüm sınıflandırma (node classification) gibi görevler üzerinden kestirimci karar alma süreçlerini desteklemekte ve proaktif öneriler üretebilen bir karar zekâsı bileşenine dönüşmektedir. Bunun yanında, büyük dil modellerinin kurumsal uygulamalarda karşılaşılan halüsinasyon riskini sınırlamak amacıyla GraphRAG yaklaşımı ele alınmakta; metinden çizge çıkarımı, topluluk temelli özetleme ve alt çizge odaklı bilgi getirimi sayesinde LLM çıktılarının izlenebilir kanıtlara dayandırılabilmesi gösterilmektedir. Sonuç olarak bölüm, referans bir mimari, karar kalitesi ölçümleme çerçevesi, yönetim ilkeleri ve geleceğe dönük araştırma gündemi önerileri üzerinden, YBS literatüründe karar zekâsı ile bilgi çizgesi entegrasyonuna bütüncül bir katkı sunmayı amaçlamaktadır.

## **2. KAVRAMSAL ÇERÇEVE: KARAR DESTEKTEN KURUMSAL KARAR ZEKÂSINA**

Karar Destek Sistemleri (KDS), bilgi sistemleri literatüründe uzun yıllardır karar vericilere veri, model ve analitik araçlar aracılığıyla destek sağlayan etkileşimli sistemler olarak tanımlanmaktadır. Zaman içinde iş zekâsı (BI), analitik ve büyük veri uygulamalarının bu yapıya eklenmesi, KDS ekosistemini hem genişletmiş hem de kavramsal sınırlarını daha geçirgen hâle

getirmiştir. Bu dönüşüm, literatürde farklı kavramların zamanla örtüşmesine ve terminolojik belirsizliklerin artmasına yol açmıştır. Yönetim Bilişim Sistemleri açısından bakıldığında ise temel eğilim açıktır (Şekil 1): Alan, yalnızca raporlama ve gösterge panoları üretmeye odaklanan geleneksel yaklaşımdan uzaklaşmakta; karar süreçlerinin iş akışlarına entegre edildiği ve karar kalitesinin sistematik biçimde yönetildiği daha bütüncül bir yapıya yönelmektedir (Phillips-Wren ve ark., 2021).



Evrimin ana yönü: raporlama odaklı destekten karar kalitesini yöneten, ölçen ve öğrenen sistemlere geçiş

## Şekil 2. Karar destek sistemlerinden kurumsal karar zekasına

Bu çerçeve, Yönetim Bilişim Sistemleri açısından iki temel hususu öne çıkarmaktadır. İlk olarak, karar zekâsı yalnızca modellerin teknik başarımı ya da tahmin doğruluğu üzerinden değerlendirilemez. Kararın üretildiği örgütsel bağlam; yetki dağılımı, süreç tasarımı, risk yapısı ve uyumluluk gereklilikleriyle birlikte ele alındığında gerçek karar kalitesi ortaya çıkmaktadır. Dolayısıyla burada söz konusu olan, yalnızca daha doğru tahmin üretmek değil, bu tahminlerin kurumsal yapı içinde nasıl anlam kazandığını ve nasıl uygulanabildiğini de yönetmektir. İkinci olarak, karar zekâsının etkin işlemesi için “bilgi” katmanını en az veri kadar kurucu bir rol oynamaktadır. Yalnızca sayısal veriler değil; politika metinleri, sözleşmeler, müşteri etkileşim kayıtları, süreç dokümantasyonu ve kurumsal hafıza unsurları da karar süreçlerine bağlam kazandıran kritik bileşenlerdir. Nitekim Gartner’ın 2024 yılı veri ve analitik eğilimlerinde karar zekâsı uygulamalarının güven, süreç ve sonuç izleme ile yönetim boyutlarıyla birlikte değerlendirilmesi, bu yaklaşımın salt teknik bir analitik çerçeve değil, aynı zamanda

güçlü bir yönetim pratiği olduğunu göstermektedir (Gartner, 2024).

Bu bağlamda, Kurumsal Karar Zekâsı ile bilgi çizgelerinin bir araya gelmesi güçlü ve doğal bir tamamlayıcılık ortaya çıkarmaktadır. Çünkü karar süreçleri bu yapıda yalnızca tablolar arasında kurulan teknik eşleştirmelere indirgenmemekte; kimlik, ilişki, bağlam ve anlamsal bütünlük üzerinden şekillenen çok kaynaklı bir bilgi ekosisteminden beslenmektedir. Başka bir ifadeyle, bilgi çizgeleri karar verme sürecine yalnızca veri entegrasyonu değil, aynı zamanda kurumsal gerçekliğin daha bütüncül ve anlamlı bir temsiline dayanan bir yorumlama kapasitesi kazandırmaktadır.

### **3. KURUMSAL BİLGİ ÇİZGELERİ: ANLAMSAL OMURGA, VERİ BÜTÜNLEŞTİRME VE YÖNETİŞİM**

Bilgi çizgeleri (knowledge graphs), literatürde gerçek dünyaya ilişkin bilgiyi depolamayı, düzenlemeyi ve aktarmayı amaçlayan çizge temelli veri yapıları olarak tanımlanmaktadır. Özellikle büyük ölçekli, dinamik ve heterojen veri kümelerinin bütünleştirilmesi ve bu verilerden anlamlı değer üretilmesi gereksinimi, bilgi çizgelerini hem akademik araştırmalarda hem de endüstriyel uygulamalarda giderek daha görünür hâle getirmiştir (Christou ve Shimizu, 2026). Bu bağlamda kurumsal bilgi çizgesi (enterprise knowledge graph), kurum bünyesinde yer alan müşteri, ürün, sözleşme, süreç, risk, kontrol, olay, çalışan ve tedarikçi gibi varlıkları ve bu varlıklar arasındaki ilişkileri anlamsal açıdan tutarlı bir yapı içinde bir araya getiren bir bilgi katmanı olarak değerlendirilebilir. Böyle bir yapı, yalnızca veri bütünleştirme işlevi görmekle kalmayıp, aynı zamanda karar süreçlerini destekleyen kurumsal bağlamın görünür ve yönetilebilir hâle gelmesine de katkı sağlamaktadır.

Teknik açıdan değerlendirildiğinde, bilgi çizgelerinin en yaygın dayanaklarından biri RDF veri modelidir. World Wide Web Consortium tarafından geliştirilen bu model, bilgiyi özne–yüklem–nesne (subject–predicate–object) biçimindeki üçlü yapılarla temsil etmektedir. Bu üçlülerin bir araya gelmesiyle oluşan yapı “RDF graph” olarak adlandırılmakta ve böylece düğüm–ilişki–düğüm biçimindeki çizgesel gösterim doğrudan desteklenmektedir (Wang ve ark., 2026). RDF’nin üzerine inşa edilen SPARQL ise bu yapılar için geliştirilen standart sorgulama dilidir. SPARQL’nin temel gücü, verinin doğrudan RDF biçiminde tutulduğu ya da bir ara katman aracılığıyla RDF olarak yorumlanabildiği farklı kaynaklar üzerinde birleşik sorgulama yapabilmesidir. Bunun yanında OWL 2, sınıflar, özellikler ve bireyler üzerinden daha zengin ontolojik modellemeye imkân tanıyan bir standart ailesi sunmaktadır (Zao ve ark., 2023). SHACL ise RDF çizgelerinin belirli kurallara ve kısıtlara göre doğrulanmasını sağlayan bir çerçeve olarak öne çıkmaktadır. Özellikle kurumsal bilgi çizgelerinde veri kalitesinin korunması, veri sözleşmelerinin tanımlanması ve uyumluluk kontrollerinin yürütülmesi bakımından SHACL önemli bir işlev üstlenmektedir.

YBS perspektifinde “kurumsal bilgi çizgesi”nin asıl gücü, doğrudan karar zekâsı ihtiyaçlarına temas eden dört yetenekte toplanabilir:

Birincisi, bilgi çizgeleri kurumsal ortamlarda kimlik sürekliliğinin sağlanması ve varlıklara ilişkin 360 derece görünürlük oluşturulması açısından kritik bir işlev görmektedir. Kurum içindeki farklı bilgi sistemlerinde aynı gerçek dünya varlığı—örneğin bir müşteri—çoğu zaman farklı anahtarlar, farklı kayıt yapıları ve farklı bağlamsal tanımlarla temsil edilmektedir. Bu durum, verinin parçalanmasına ve karar süreçlerinde aynı varlığın birden fazla, hatta çelişkili biçimde ele alınmasına yol açabilmektedir. Bilgi çizgeleri ise kimlik birleştirme ve ilişki haritalama mekanizmaları sayesinde bu

parçalı yapıyı daha bütüncül bir kurumsal temsile dönüştürmektedir. Böylece karar süreçleri yanlış, eksik ya da yinelenmiş kayıtlar yerine daha doğru tanımlanmış varlıklar üzerinden yürütülebilmektedir. Aksi durumda karar zekâsı sistemleri, tutarsız veya mükerrer varlık kayıtlarından beslenerek hatalı optimizasyonlara ve yanıltıcı önerilere yol açabilmektedir (Christou ve Shimizu, 2026).

İkincisi, bilgi çizgeleri bağlamsal entegrasyon kapasitesi sayesinde kurumsal karar süreçlerine önemli bir derinlik kazandırmaktadır. Bu yapılar yalnızca veri satırlarını ya da izole kayıtları bir araya getirmekle kalmaz; süreçleri, kuralları, risk unsurlarını, kontrol mekanizmalarını ve ilgili belgeleri de ilişkisel bir bağlam içinde bütünleştirir. Böylece karar verme süreci, yalnızca gerçekleşen olayların betimlenmesine indirgenmez; olayların hangi nedenlerle ortaya çıktığı ve benzer koşullar altında tekrar edip etmeyeceği gibi daha açıklayıcı ve öngörüye dayalı soruların da ele alınmasına imkân tanır. Bu yönüyle bilgi çizgeleri, karar destek kapasitesini betimleyici analizden açıklayıcı ve kestirimci analize doğru genişleten bir altyapı sunmaktadır (Hogan ve ark., 2021).

Üçüncüsü, bilgi çizgeleri karar süreçlerinde izlenebilirlik ve denetlenebilirlik gereksinimlerini karşılamada önemli bir avantaj sağlamaktadır. Özellikle finans, sağlık ve kamu gibi yüksek düzeyde düzenlemeye tabi sektörlerde, bir kararın hangi bilgiye, hangi kurala ve hangi ilişki ağına dayanarak üretildiğinin açıklanabilmesi kritik bir beklenti hâline gelmiştir. Bu bağlamda bilgi çizgeleri, karar girdilerinin kökenini ve dönüşüm sürecini kaynak-soy ağacı (lineage) mantığıyla takip etmeye elverişli bir temsil sunmaktadır. Böylece kararların dayandığı kanıt zinciri daha görünür, daha denetlenebilir ve gerektiğinde geriye dönük olarak doğrulanabilir hâle gelmektedir. Bu özellik, yalnızca teknik şeffaflığı artırmakla kalmayıp, aynı zamanda kurumsal

güven, uyumluluk ve hesap verebilirlik açısından da önemli bir katkı sunmaktadır.

Dördüncüsü, bilgi çizgelerinin kurumsal değer üretebilmesi, yalnızca başlangıçta doğru biçimde modellenmesine değil, aynı zamanda veri ve bilgi yaşam döngüsünün etkin biçimde yönetilmesine bağlıdır. Bir bilgi çizgesinin sürdürülebilirliği; şema ve ontoloji yönetimi, veri kalite kurallarının tanımlanması ve doğrulanması, erişim yetkilerinin düzenlenmesi, değişikliklerin kontrollü biçimde yönetilmesi ve sistemin sürekli gözlemlenebilir olması gibi yönetim unsurlarıyla doğrudan ilişkilidir. Bu nedenle bilgi çizgesi, statik bir veri modeli olarak değil, bakım, güncelleme ve kalite güvencesi gerektiren yaşayan bir kurumsal varlık olarak ele alınmalıdır. Özellikle kurumsal karar zekâsı bağlamında, kararların güvenilirliği ve tutarlılığı büyük ölçüde bilgi çizgesinin kalitesine bağlı olduğundan, KG kalitesinin de ölçülebilir ve yönetilebilir bir unsur olarak değerlendirilmesi gerekmektedir (Ye ve ark., 2026).

Literatürde bilgi çizgelerinin Karar Destek Sistemleri tasarımına entegrasyonu da açık biçimde ele alınmıştır. Elnagar ve Weistroffer, KDS tasarımında bilgi edinimi ve bilginin temsiliyle ilgili temel sorunların bilgi çizgeleri aracılığıyla daha esnek biçimde ele alınabileceğini ileri sürmektedir. Yazarlara göre, ontoloji temelli temsiller güçlü bir anlamsal çerçeve sunsa da gerçek zamanlı güncelleme ve dinamik bilgi işleme gereksinimleri karşısında bazı sınırlılıklar taşımaktadır; buna karşılık bilgi çizgeleri daha çevik ve güncellenebilir bir alternatif sunabilmektedir (Elnagar ve Weistroffer, 2019). Bu yaklaşım, bilgi çizgeleri ile karar zekâsının birlikte ele alınmasının Yönetim Bilişim Sistemleri literatüründeki kuramsal dayanağını daha da güçlendirmektedir. Bu çerçevede bilgi çizgeleri, karar kalitesini artırmak amacıyla kurumsal bilgiye erişimi kolaylaştıran, bilgiyi bağlam içinde kullanılabilir hâle getiren ve karar süreçlerine

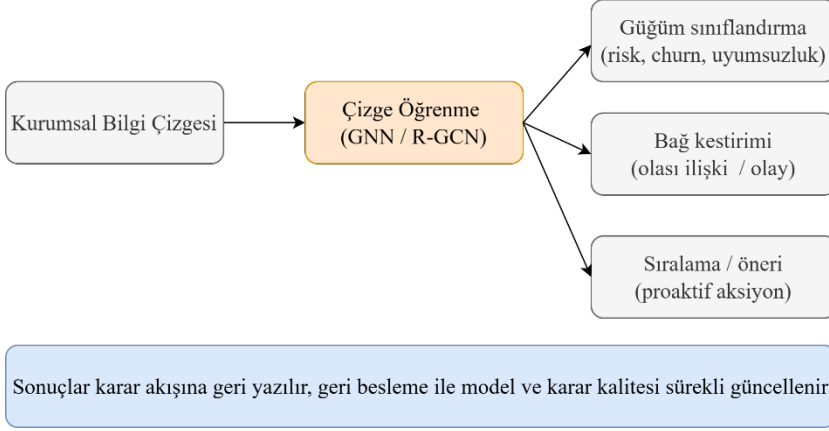
anlamsal derinlik kazandıran temel bir altyapı katmanı olarak konumlanmaktadır.

#### **4. GNN İLE KESTİRİMCİ KARAR ALMA: KURUMSAL BİLGİ ÇİZGELERİ ÜZERİNDE ÇİZGE ÖĞRENME**

Kurumsal karar zekâsı, birçok durumda yalnızca mevcut durumu görünür kılmayı değil, aynı zamanda olası gelecek senaryolarını öngörmeyi ve bu senaryolar karşısında en uygun eylem seçeneklerini belirlemeyi amaçlamaktadır. Bilgi çizgeleri bu bağlamda iki yönlü bir değer üretmektedir. İlk olarak, müşteriler, ürünler, süreçler, riskler, sözleşmeler ve kanallar gibi kurumsal yapının yüksek derecede ilişkisel bileşenlerini doğal biçimde temsil edebilen bir anlamsal çerçeve sunmaktadır. İkinci olarak ise bu temsil, çizge öğrenme yöntemleriyle işlendiğinde kestirimci sinyaller üretebilen analitik bir zemine dönüşmektedir. Böylece bilgi çizgeleri, yalnızca kurumsal gerçekliğin modellenmesini değil, aynı zamanda bu gerçeklik üzerinden geleceğe dönük karar desteği üretilmesini de mümkün kılmaktadır.

Çizge sinir ağları (GNN), düzensiz topolojik yapıya sahip çizge verileri üzerinde anlamlı gösterimler (embeddings) öğrenmek amacıyla geliştirilen bir derin öğrenme yaklaşımıdır. Geleneksel derin öğrenme yöntemlerinin çoğu düzenli veri yapıları üzerinde daha etkili çalışırken, GNN mimarileri düğümler, kenarlar ve bunlar arasındaki karmaşık ilişkilerden öğrenebilecek biçimde tasarlanmıştır. Bu literatürde özellikle mesaj geçişi (message passing), komşuluk bilgisinin toplanması ve birleştirilmesi (aggregation) ile çok ilişkili çizge yapılarının modellenmesi temel tartışma eksenlerini oluşturmaktadır (Şekil 3). GNN alanını bütüncül biçimde ele alan güncel derleme çalışmaları da yöntemsel sınıflandırmaları, temel mimari

farklılıkları ve uygulama alanlarını ayrıntılı biçimde incelemektedir (Shabani ve ark., 2024).



**Şekil 3. GNN ile kestirimci karar alma**

Yönetim Bilişim Sistemleri bağlamında kurgulanmış bir kurumsal bilgi çizgesi üzerinde GNN kullanımının en önemli avantajlarından biri, departmanlar arası ya da çapraz alanlı örüntüleri görünür hâle getirebilmesidir. Çünkü kurumsal gerçeklik çoğu zaman tek bir veri tablosu ya da tekil bir işlevsel sistem üzerinden tam olarak temsil edilemez. Satış, tedarik, finans, risk yönetimi, bilgi teknolojileri hizmet yönetimi ve uyumluluk gibi farklı alanlara ait veriler birlikte değerlendirildiğinde, kurumun işleyişine ilişkin daha bütüncül bir görünüm elde edilir. Bu bütüncül yapı, yalnızca veri entegrasyonunu güçlendirmekle kalmaz, aynı zamanda karar süreçlerinde gözden kaçabilecek ilişkileri ortaya çıkararak karar kalitesinin artırılmasına da katkı sağlar.

Kurumsal bilgi çizgileri üzerinde ele alınan karar problemleri çoğu zaman üç temel görev etrafında yapılandırılabilir. Bunlardan ilki düğüm sınıflandırmadır (node classification). Bu tür görevlerde, örneğin bir tedarikçinin risk düzeyinin belirlenmesi, bir müşterinin terk etme olasılığının sınıflandırılması ya da bir iş sürecinin uyumsuzluk riski

bakımından değerlendirilmesi hedeflenmektedir. İkinci temel görev bağlantı kestirimidir (link prediction). Bu kapsamda müşteri ile ürün arasındaki olası satın alma ilişkisi, tedarikçi ile gecikme olayı arasındaki muhtemel nedensel bağ veya bir BT olayı ile onun kök nedeni arasındaki ilişkinin tahmin edilmesi söz konusu olmaktadır. Üçüncü görev ise sıralama ve öneri üretimidir (ranking/recommendation). Bu aşamada amaç, hangi tedarikçinin devreye alınmasının daha uygun olacağı, hangi aksiyon planının riski daha fazla azaltacağı ya da hangi müdahale seçeneğinin daha yüksek kurumsal fayda sağlayacağı gibi alternatifleri önceliklendirmektir. Bu üçlü yapı, kurumsal karar zekâsı uygulamalarında bilgi çizgesi tabanlı analitiğin hem açıklayıcı hem de eyleme dönük boyutunu görünür kılmaktadır.

Bu görevler doğrultusunda literatürde öne çıkan temel mimari hatlardan biri, Graph Convolutional Networks (GCN) yaklaşımıdır. GCN temelli modeller, özellikle büyük ölçekli çizgelerde temsil öğrenimini mümkün kılan ve zamanla daha ölçeklenebilir varyantlarla genişleyen bir gelişim çizgisi ortaya koymuştur. Bununla birlikte kurumsal bilgi çizgeleri çoğu durumda yalnızca basit bağlantılardan oluşan yapılar değildir; aksine, aynı düğüm türleri arasında farklı ilişki tiplerini, yönlü bağlantıları ve farklı varlık türleri arasındaki çok katmanlı etkileşimleri içeren çok ilişkili (multi-relational) yapılardır. Bu nedenle kurumsal bağlamda yalnızca klasik çizge evrişimli yaklaşımlar çoğu zaman yeterli olmamakta, ilişki türlerini açık biçimde modele dâhil edebilen daha gelişmiş mimarilere ihtiyaç duyulmaktadır. Bu gereksinime yanıt veren yaklaşımlardan biri olan Relational GCN (R-GCN), çok ilişkili çizgeler üzerinde öğrenme yapabilme kapasitesi sayesinde özellikle bilgi çizgesi tamamlama, bağlantı kestirimi ve varlık sınıflandırma gibi görevlerde doğrudan kullanılmaktadır (Fu ve ark., 2021).

Teknik derinliği somutlaştırmak açısından, kurumsal bilgi çizgesi üzerinde GNN kullanımına ilişkin örnek bir model

senaryosu da önerilebilir. Örneğin kurumsal risk ve süreç bileşenlerini içeren bir bilgi çizgesi üzerinde bağlantı kestirimi (link prediction) kullanılarak “kontrol zafiyeti → olası olay” ilişkileri önceden tahmin edilebilir. Bu yapı içerisinde süreç adımları, üçüncü taraf tedarikçiler veya BT varlıkları gibi kırılgan düğümler hem merkezilik ölçütleri hem de GNN tabanlı risk skorları yardımıyla sıralanabilir. Elde edilen sonuçlar doğrultusunda, hangi kontrol mekanizmalarının öncelikle güçlendirilmesi gerektiği, hangi tedarikçilerin daha yakından izlenmesi gerektiği ya da hangi BT varlıklarının yüksek risk taşıdığına ilişkin proaktif aksiyon önerileri üretilebilir. Bu tür bir senaryo, GNN’nin teknik kapasitesi ile işletme açısından anlamlı sonuçları—özellikle risk azaltımı, uyumluluk yönetimi ve karar önceliklendirmesi—aynı çerçevede buluşturduğu için YBS odaklı bir katkıyı oldukça somut biçimde görünür kılmaktadır.

## **5. KG VE LLM ENTEGRASYONU: GRAPHRAG İLE GÜVENİLİR KURUMSAL ASİSTANLAR**

Kurumsal Karar Zekâsı’nın 2024–2026 dönemindeki en dikkat çekici gelişim eksenlerinden biri, büyük dil modelleriyle desteklenen konuşma tabanlı analitik ve karar asistanlarının yaygınlaşmasıdır. Doğal dil üzerinden sorgulama, özetleme, açıklama ve öneri üretme kapasitesi, bu sistemleri kurumsal karar süreçleri açısından oldukça cazip hale getirmiştir. Bununla birlikte, LLM’lerin kurumsal ortamlardaki kullanımı güvenilirlik, doğruluk ve özellikle halüsinasyon olarak adlandırılan gerçeğe aykırı içerik üretimi riski nedeniyle önemli sınırlılıklar taşımaktadır. LLM halüsinasyonlarını sınıflandıran ve nedenleri ile azaltım stratejilerini tartışan kapsamlı çalışmalar, bu problemin veri, eğitim ve çıkarım aşamalarının her birinde ortaya çıkabildiğini; kurumsal bağlamda ise güven, açıklanabilirlik ve

denetlenebilirlik gereksinimleri nedeniyle çok daha kritik bir nitelik kazandığını göstermektedir (Huang ve ark., 2025).

Bu riskleri azaltmak amacıyla geliştirilen temel yaklaşımlardan biri Retrieval-Augmented Generation (RAG) çerçevesidir. RAG yaklaşımı, modelin yanıt üretmeden önce harici bir bilgi kaynağından ya da doküman koleksiyonundan ilgili içerikleri getirmesini ve böylece yanıt yalnızca parametrelerine gömülü bilgiye değil, erişilebilir bir bağlamsal dayanağa yaslamasını amaçlamaktadır. Bu yaklaşımın temel çalışmalarından biri, bilgi yoğun görevlerde getirici (retriever) ve üretici (generator) bileşenlerin birlikte çalıştığı hibrit bir yapı önererek parametrelili ve parametresiz belleğin nasıl bütünleştirilebileceğini göstermiştir (Lewis ve ark., 2020). Böylece amaç, model çıktılarının daha güncel, daha ilgili ve daha kanıt temelli hale getirilmesidir.

Bununla birlikte, klasik RAG mimarileri kurumsal uygulamalarda iki temel sorunla karşılaşmaktadır. İlk sorun, doğru içerik parçasının bulunması ve uygun bağlamın seçilmesiyle ilgilidir. Getirim doğruluğu, parçalama stratejileri ve bağlam penceresinin verimli kullanımı bu noktada belirleyici hale gelmektedir. İkinci sorun ise farklı belge parçaları arasında anlamlı ilişkiler kurabilme kapasitesidir. Pek çok kurumsal karar problemi, yalnızca tek bir dokümana ya da tek bir kayıt parçasına dayanarak çözülememekte; farklı kaynaklardan gelen bilgilerin ilişkilendirilmesini gerektirmektedir. Bilgi çizgeleri tam da bu aşamada RAG için yapısal bağlam sağlayan güçlü bir çözüm sunmaktadır. Çünkü bilgi çizgeleri, doküman parçalarını, varlıkları, olayları ve süreçleri ilişkisel bir bütünlük içinde bağlayarak getirmenin yalnızca tekil metin parçaları üzerinden değil, anlamlı alt çizgeler üzerinden yapılmasına olanak tanımaktadır. Bu yaklaşım, kurumsal gerçekliğin daha tutarlı ve bağlamsal biçimde temsil edilmesini mümkün kılmaktadır.

Bu genel düşünceyi daha sistematik ve kurumsal ölçekte ele alan dikkat çekici yaklaşımlardan biri GraphRAG'dir. Microsoft Research tarafından geliştirilen bu yaklaşım, metinden varlık ve ilişki çıkarımı, ağ analizi ve LLM tabanlı özetleme/prompt tasarımı bileşenlerini bir araya getirerek özellikle anlatı niteliği taşıyan özel metin koleksiyonları üzerinde keşif, özetleme ve soru-cevap süreçleri için uçtan uca bir yöntem sunmaktadır. GraphRAG yaklaşımının temel iddiası, özel veri kümeleri üzerinde çizge tabanlı bir bilgi organizasyonu sayesinde hem daha geniş ölçekli sorguların desteklenebilmesi hem de daha tutarlı ve ilişkisel açıdan zengin bağlamların üretilebilmesidir. Bu yönüyle GraphRAG, yalnızca doküman getirimine dayalı klasik RAG yaklaşımlarına göre daha yüksek yapısal farkındalık sunmaktadır.

Akademik ve kurumsal açıdan daha güçlü bir katkı üretmek için GraphRAG yaklaşımını doğrudan kurumsal bilgi çizgesi kavramıyla ilişkilendirmek önem taşımaktadır. Kurumsal bilgi çizgesi yalnızca dokümanlardan çıkarılan varlık ve ilişkilerden ibaret değildir; aynı zamanda ERP ve CRM süreçleri, iş kuralları, risk kontrolleri, veri sözlüğü, politika belgeleri ve yönetsel metaveri gibi unsurları da kapsamaktadır. Bu geniş çerçevede, RAG sistemlerinin yalnızca doküman düzeyinde değil, kurumsal gerçekliğin bütüncül yapısı içinde tutarlı yanıtlar üretmesine katkı sağlamaktadır. Benzer biçimde GraphRAG'ın topluluk temelli özetleme yaklaşımı, kurumsal karar bağlamında departman, ürün ailesi, olay kümesi ya da süreç segmenti gibi karar merkezli birimlerin özetlenmesi için uyarlanabilir. Böylece sistem yalnızca tekil sorgulara yanıt veren bir araç olmaktan çıkarak, kurum içindeki ilişkisel yapıları karar verici için anlamlı ve özetlenmiş bilgi kümelerine dönüştürebilir. Ayrıca bilgi çizgesi tabanlı RAG varyantlarının parça genişletme, ilişkisel örgütlenme ve bağlamsal çeşitlilik artırma gibi mekanizmalarla getirim sürecini iyileştirmeye çalıştığı gösterilmiştir. Bu

yaklaşım, kurumsal karar zekâsında tek bir dokümana sıkışmayan, çok kaynaklı ve ilişki temelli kanıt üretme gereksinimiyle doğrudan örtüşmektedir.

## **6. UÇTAN UCA REFERANS MİMARİ: KURUMSAL KG OMURGALI KARAR ZEKÂSI**

Bu bölümün teknik açıdan güçlü, ancak aynı zamanda Yönetim Bilişim Sistemleri eksenine sadık bir çerçeve içinde kalabilmesi için en etkili yaklaşım, tartışılan kavramların bütünleşik bir referans mimari altında bir araya getirilmesidir. Böyle bir mimari; karar zekâsının karar-merkezli yaklaşımını, kurumsal bilgi çizgelerinin anlamsal omurga işlevini ve GNN ile GraphRAG katmanlarının analitik kapasitesini tek bir sistematik yapı içinde buluşturmaktadır.

Bu yapının başlangıç noktasını veri kaynakları oluşturmaktadır. Kurumsal kararların girdileri yalnızca veri ambarlarında tutulan yapılandırılmış tablolardan ibaret değildir. Süreç kayıtları, olay logları, dokümanlar, e-postalar, toplantı metinleri, politika belgeleri ve prosedürler de karar süreçleri açısından kritik öneme sahip bilgi kaynaklarıdır. Bu veri çeşitliliği, bilgi çizgelerinin ortaya çıkış mantığıyla doğrudan uyumludur; çünkü bilgi çizgeleri tam da bu tür heterojen ve dağınık bilgi kümelerini ortak bir anlamsal yapı altında bütünleştirmek için güçlü bir zemin sunmaktadır (Hogan ve ark., 2021).

İkinci katman, kurumsal bilgi çizgesinin üretim hattını oluşturmaktadır. Bu aşamada öncelikle kavramsal şema ve ontoloji tasarımı yoluyla ortak bir kurumsal dil inşa edilmektedir. OWL tabanlı ontolojik modelleme, kurum içindeki varlıkların, kavramların ve ilişkilerin tutarlı bir biçimde tanımlanmasına olanak sağlamaktadır. Bunu izleyen aşamada RDF temelli veri dönüştürme ve ilişkilendirme süreçleri sayesinde farklı

kaynaklarda dağınık biçimde bulunan varlıklar ortak bir çizgesel yapı içinde bağlanmaktadır. Üçüncü adımda SHACL kullanılarak kalite ve doğrulama mekanizmaları devreye alınmakta, böylece bilgi çizgesinin karar süreçleri için güvenilir bir bilgi katmanı hâline gelmesi sağlanmaktadır. Son olarak SPARQL, hem analitik amaçlı sorgulama hem de operasyonel erişim gereksinimleri için ortak bir erişim dili sunmaktadır. Bu katman aynı zamanda kararların hangi bilgi kaynaklarına dayanılarak üretildiğini izlemeyi mümkün kılan kaynak ve soy ağacı yönetimi açısından da önem taşımaktadır. Dolayısıyla bu yapı yalnızca veri bütünleştirme değil, aynı zamanda karar yönetişimi için de temel oluşturmaktadır.

Üçüncü katman, çizge analitiği ve çizge öğrenme bileşenlerinden oluşmaktadır. Bu aşamada GNN tabanlı yöntemler yalnızca kestirim üreten teknik araçlar olarak değil, karar akışının etkin bir parçası olarak ele alınmaktadır. Başka bir ifadeyle GNN, kurumsal bilgi çizgesi üzerinde örüntüleri ortaya çıkaran, riskli ya da fırsat içeren ilişkileri tahmin eden ve karar vericilere proaktif öneriler sunan bir karar destek bileşenine dönüşmektedir. GNN'lerin yöntemsel çerçevesi ve çok ilişkili çizgelere nasıl uyarlanabildiği literatürde geniş biçimde tartışılmaktadır. Özellikle bilgi çizgelerinde eksik bağlantıların tamamlanması, bağlantı olasılıklarının kestirilmesi ve düğüm özelliklerinin sınıflandırılması gibi görevler, kurumsal bağlamda doğrudan karar problemlerine çevrilebilmektedir (Zhou ve ark., 2020). Bu yönüyle çizge öğrenme katmanı, bilgi çizgesini pasif bir temsil yapısından aktif bir karar zekâsı altyapısına dönüştürmektedir.

Dördüncü katman, büyük dil modelleriyle desteklenen karar asistanı bileşenidir. Burada klasik RAG yaklaşımından GraphRAG çizgisine uzanan yapı öne çıkmaktadır. Temel amaç, getirmenin yalnızca doküman parçaları düzeyinde değil, alt çizge temelli kanıt kümeleri üzerinden yapılmasıdır. Böylece kurumsal

karar asistanı, yalnızca ilgili metin parçalarını geri çağıran bir sistem olmaktan çıkarak, karar bağlamını ilişkisel ve anlamsal bir bütünlük içinde sunabilen bir yapıya kavuşmaktadır. GraphRAG'in metin çıkarımı, çizge üretimi ve çizge temelli özetleme yaklaşımı, özellikle açıklanabilirlik ve izlenebilirlik gerektiren kurumsal karar ortamlarında önemli avantajlar sağlamaktadır (Sesen ve ark., 2025). Bu sayede LLM çıktıları, daha güçlü bağlamsal destek ve daha görünür kanıt yapıları üzerine inşa edilebilmektedir.

Beşinci katman ise karar tasarımı ve yürütüm boyutunu kapsamaktadır. Bu aşama, sistemin YBS karakterini en belirgin biçimde ortaya koyan katmandır. Çünkü burada yalnızca analitik çıktı üretmek değil, karar modelinin açık biçimde tanımlanması, karar akışlarının iş süreçlerine gömülmesi, insan onayı gerektiren adımların belirlenmesi, KPI ve OKR yapılarıyla ilişki kurulması ve karar kalitesinin sistematik biçimde ölçülmesi gibi yönetsel bileşenler devreye girmektedir. Böylece karar zekâsı, yalnızca teknik olarak başarılı tahminler üreten bir sistem olmaktan çıkarak, kurumsal hedeflerle uyumlu, ölçülebilir ve yönetilebilir bir karar altyapısı hâline gelmektedir. Bu yaklaşım, karar zekâsı platformlarının kararları açık biçimde modelleme, akışı yürütme, karar kalitesini izleme ve sonuçlardan öğrenme yönündeki temel vurgusuyla doğrudan örtüşmektedir (Adigopula ve Siriki, 2025).

Bu mimarinin kurumsal uygulamadaki karşılığını göstermek açısından endüstriden verilecek kısa ancak güçlü örnekler bölümün inandırıcılığını artıracaktır. Örneğin Netflix tarafından paylaşılan Unified Data Architecture yaklaşımı, kavramsal alan modellerini tutarlı şemalara ve veri boru hatlarına dönüştüren, bilgi çizgesi temelli bir kurumsal mimari örneği olarak değerlendirilebilir. Bu örnek, bilgi çizgelerinin yalnızca teorik bir modelleme aracı olmadığını, aynı zamanda büyük ölçekli kurumsal mimarilerin çekirdeğinde yer alabilecek pratik bir altyapı sunduğunu göstermektedir (Jiang ve ark., 2025).

## **7. ARAŞTIRMA GÜNDEMİ VE SONUÇ**

Mevcut çalışmalar incelendiğinde, bilgi çizgelerinin çoğunlukla veri bütünleştirme, semantik modelleme ve bilgi yönetimi ekseninde ele alındığı; karar zekâsı yaklaşımının ise daha çok karar süreçlerinin mühendisliği, platformlaşma ve yönetim boyutları üzerinden tartışıldığı görülmektedir. Her ne kadar Karar Destek Sistemleri tasarımında bilgi çizgelerinin kullanımına değinen çalışmalar bulunsa da, kurumsal karar zekâsının karar-merkezli geri besleme mantığını, kurumsal bilgi çizgelerini, GNN tabanlı kestirimci analitiği ve GraphRAG destekli güvenilir kurumsal asistan yaklaşımını aynı çatı altında birleştiren bütüncül çerçeveler hâlen oldukça sınırlıdır. Bu nedenle söz konusu yaklaşım, yalnızca teknik bileşenleri yan yana getiren bir öneri değil, aynı zamanda Yönetim Bilişim Sistemleri disiplinde karar verme süreçlerini yeniden düşünmeye imkân tanıyan bütünleşik bir kuramsal zemin sunmaktadır.

özgün katkıların üç başlık altında netleştirilmesi oldukça etkili olacaktır. İlk katkı, karar-merkezli anlamsal omurga yaklaşımıdır. Burada bilgi çizgesi, kurum içinde yalnızca bir veri entegrasyon katmanı olarak değil, aynı zamanda kararların hangi bilgiye dayandığını görünür kılan, izlenebilirlik sağlayan ve karar kalitesinin yönetimine hizmet eden bir anlamsal katman olarak konumlandırılmaktadır. Bu yönüyle bilgi çizgesi, karar kanıt zincirinin kurulmasına ve kararların kurumsal hafıza içinde denetlenebilir hale gelmesine katkı sunmaktadır (Elnagar ve Weistroffer, 2019).

İkinci katkı, kestirimci ve proaktif karar desteği boyutunda ortaya çıkmaktadır. Kurumsal bilgi çizgesi üzerinde çalışan GNN tabanlı modeller, departmanlar arası ya da çapraz alanlı örüntüleri görünür hâle getirerek erken uyarı, öneri ve önceliklendirme mekanizmalarını destekleyebilmektedir. Risk

tahmini, müşteri terk olasılığı, tedarik gecikmeleri veya süreç kırılganlıkları gibi konularda üretilecek sinyaller, karar vericilere yalnızca geçmişini açıklayan değil, aynı zamanda geleceğe dönük eylem önerileri sunan bir yapı kazandırmaktadır. Böylece bilgi çizgesi, pasif bir temsil yapısı olmaktan çıkıp proaktif kurumsal zekânın analitik zemini haline gelmektedir.

Üçüncü katkı ise güvenilir kurumsal LLM tasarımına ilişkindir. GraphRAG yaklaşımı aracılığıyla kurumsal metinler ile bilgi çizgelerinin bir araya getirilmesi, büyük dil modellerinin ürettiği yanıtların daha güçlü bağlamsal dayanaklara ve daha görünür kanıtlara yaslanmasını mümkün kılmaktadır. Bu yaklaşım, halüsinasyon riskini tamamen ortadan kaldırdığı iddiasında bulunmaktan ziyade, bu riski ölçümleme, izleme ve yönetim mekanizmalarıyla birlikte yönetilebilir hale getirmeyi hedeflemektedir. Böylece kurumsal karar asistanları, yalnızca doğal dilde yanıt veren araçlar olmaktan çıkıp, açıklanabilir ve kurumsal gerçeklikle daha uyumlu karar destek bileşenlerine dönüşmektedir.

## KAYNAKÇA

- Adigopula, D., & Siriki, G. (2025). *Next-generation performance KPI systems: Integrating predictive analytics and OKRs* (Version 2). Zenodo. <https://doi.org/10.5281/zenodo.17654110>
- Christou, A., & Shimizu, C. (2026). Experiments in graph structure and knowledge graph embeddings. *Neurosymbolic Artificial Intelligence*, 2. <https://doi.org/10.1177/29498732261420038>
- Elnagar, S., & Weistroffer, H. R. (2019). Introducing knowledge graphs to decision support systems design. In S. Wrycza & J. Maślankowski (Eds.), *Information systems: Research, development, applications, education: SIGSAND/PLAIS 2019* (Lecture Notes in Business Information Processing, Vol. 359, pp. [sayfa aralığı]). Springer. [https://doi.org/10.1007/978-3-030-29608-7\\_1](https://doi.org/10.1007/978-3-030-29608-7_1)
- Fu, S., Liu, W., Zhang, K., & Zhou, Y. (2021). Example-feature graph convolutional networks for semi-supervised classification. *Neurocomputing*, 461, 63–76. <https://doi.org/10.1016/j.neucom.2021.07.048>
- Gartner. (2024). *Gartner identifies the top trends in data and analytics for 2024* [Basın bülteni]. Erişim tarihi: 05.03.2026, <https://www.gartner.com/en/newsroom/press-releases/2024-04-25-gartner-identifies-the-top-trends-in-data-and-analytics-for-2024>
- Hogan, A., Blomqvist, E., Cochez, M., D'Amato, C., de Melo, G., Gutierrez, C., Kirrane, S., Labra Gayo, J. E., Navigli, R., Neumaier, S., Ngomo, A.-C. N., Polleres, A., Rashid, S. M., Rula, A., Schmelzeisen, L., Sequeda, J., Staab, S., & Zimmermann, A. (2021). Knowledge graphs. *ACM*

- Computing Surveys*, 54(4), 1–37.  
<https://doi.org/10.1145/3447772>
- Huang, L., Yu, W., Ma, W., Zhong, W., Feng, Z., Wang, H., Chen, Q., Peng, W., Feng, X., Qin, B., & Liu, T. (2025). A survey on hallucination in large language models: Principles, taxonomy, challenges, and open questions. *ACM Transactions on Information Systems*, 43(2), 1–55.  
<https://doi.org/10.1145/3703155>
- Jiang, X., Wang, S., Gong, Y., Liu, L., Yu, T., & Yu, X. (2025). UDA: Unified pretraining for multiarchitecture binary disassembly. *IEEE Internet of Things Journal*, 12(14), 27291–27306.  
<https://doi.org/10.1109/JIOT.2025.3562205>
- Kostopoulos, G., Stefani, A., Vasiliadis, V., & Kotsiantis, S. (2026). Deep learning for e-commerce: Recent developments in prediction, personalization and decision intelligence. *Applied Sciences*, 16(5), 2263.  
<https://doi.org/10.3390/app16052263>
- Lewis, P., Perez, E., Piktus, A., Petroni, F., Karpukhin, V., Goyal, N., Küttler, H., Lewis, M., Yih, W., Rocktäschel, T., Riedel, S., & Kiela, D. (2020). Retrieval-augmented generation for knowledge-intensive NLP tasks. In *Advances in Neural Information Processing Systems 33 (NeurIPS 2020)*. <https://arxiv.org/abs/2005.11401>
- Phillips-Wren, G., Daly, M., & Burstein, F. (2021). Reconciling business intelligence, analytics and decision support systems: More data, deeper insight. *Decision Support Systems*, 146, 113560.  
<https://doi.org/10.1016/j.dss.2021.113560>
- Sesen, M. B., Au Yeung, J., & Asgari, E. (2025). Development and validation of retrieval augmented generation (RAG)

and GraphRAG for complex clinical cases. *medRxiv*.  
<https://doi.org/10.1101/2025.11.25.25341010>

Shabani, N., Wu, J., Beheshti, A., Sheng, Q. Z., Foo, J., Haghighi, V., Hanif, A., & Shahabikargar, M. (2024). A comprehensive survey on graph summarization with graph neural networks. *IEEE Transactions on Artificial Intelligence*, 5(8), 3780–3800.  
<https://doi.org/10.1109/TAI.2024.3350545>

Wang, J., Yang, J., Deng, J., Gunes, H., & Song, S. (2026). Graph in graph neural network. *International Journal of Computer Vision*, 134, 162.  
<https://doi.org/10.1007/s11263-026-02731-4>

Ye, J., Li, X., Nie, L., Song, K., Zhan, D., & Sun, C. (2026). Maintenance decision support system for multiple units based on high-quality knowledge graphs. *IEEE Transactions on Computational Social Systems*, XX(YY), ZZZ–ZZZ. <https://doi.org/10.1109/TCSS.2026.3653820>

Zhao, Z., Ge, X., Shen, Z., Hu, C., & Wang, H. (2023). S2CTrans: Building a bridge from SPARQL to Cypher. In C. Strauss, T. Amagasa, G. Kotsis, A. M. Tjoa, & I. Khalil (Eds.), *Database and Expert Systems Applications: DEXA 2023* (Lecture Notes in Computer Science, Vol. 14146, pp. [sayfa aralığı]). Springer. [https://doi.org/10.1007/978-3-031-39847-6\\_33](https://doi.org/10.1007/978-3-031-39847-6_33)

# **FİRMALARDA YAPAY ZEKÂ VE BÜYÜK DİL MODELLERİNİN BENİMSENMESİ: VERİ MADENCİLİĞİ TEMELLİ ÜLKELER ARASI BİR ANALİZ**

**Ahmet AYAZ<sup>1</sup>**

## **1. GİRİŞ**

Yapay zekâ (YZ) ve özellikle büyük dil modelleri (BDM), günümüzde firmaların dijital dönüşüm stratejilerinin merkezinde yer almaktadır. ChatGPT, Claude, Bard, Midjourney ve Stable Diffusion gibi üretken yapay zekâ araçları, müşteri etkileşiminden bilgi yönetimine, içerik üretiminden karar destek sistemlerine kadar çok sayıda kurumsal süreçte dönüştürücü bir rol oynamaktadır. Ancak bu teknolojilerin benimsenme düzeyleri, ülkeler, sektörler, firma büyüklükleri ve kullanım amaçları arasında dikkate değer farklılıklar göstermektedir (Shen & Li, 2025).

Bu farklılıklar, yalnızca teknolojik altyapı düzeyinden değil; aynı zamanda örgütsel hazırlık, liderlik desteği, insan sermayesi niteliği ve çevresel koşullar gibi çok boyutlu etkenlerden de kaynaklanmaktadır (Palade & Căruțașu, 2023). Bu nedenle, yapay zekâ benimsemesi yalnızca bir teknoloji adaptasyonu değil, stratejik yönetim ve dijital olgunluk süreçlerini kapsayan bir organizasyonel dönüşüm olarak ele alınmalıdır.

---

<sup>1</sup> Dr. Öğr. Üyesi, Karadeniz Teknik Üniversitesi, İktisadi ve İdari Bilimler Fakültesi, Yönetim Bilişim Sistemleri Bölümü, ORCID: 0000-0003-1405-0546.

Bu çalışmanın amacı, firmaların YZ ve BDM teknolojilerini benimseme düzeylerini veri madenciliği temelli bir keşifsel analiz aracılığıyla incelemektir. Araştırma, firmaların YZ/BDM kullanım örüntülerini nicel olarak ortaya koymakta ve bu örüntüleri ülke, sektör, firma büyüklüğü ve kullanılan yapay zekâ aracı boyutlarında karşılaştırmaktadır.

Yöntem olarak, K-Means kümeleme algoritması kullanılmış ve firmalar yüksek, orta ve düşük benimseme düzeyleri olmak üzere üç ana profile ayrılmıştır. Bu kümeleme yaklaşımı, firmaların yalnızca teknolojiyi benimseyip benimsemediklerini değil, aynı zamanda günlük kullanım yoğunluğu ve organizasyonel olgunluk düzeylerine göre farklılaştıklarını ortaya koymaktadır (Kou & Tang, 2025).

Bu yönüyle araştırma, literatürdeki üç temel boşluğu hedeflemektedir:

1. Yapay zekâ benimsemesini gerçek kullanım verileri üzerinden ölçmesi,
2. Analizi yalnızca ülke düzeyiyle sınırlamayıp sektörel, ölçeksel ve araç bazlı farklılıkları da kapsaması,
3. Veri madenciliği yöntemini kullanarak YZ/BDM olgunluk profillerini nicel biçimde keşfetmesidir.

Dolayısıyla bu çalışma, yapay zekâ ve büyük dil modellerinin küresel düzeydeki yayılımını yalnızca bir teknoloji eğilimi olarak değil, aynı zamanda firmaların dijital dönüşüm yolculuğunda ortaya çıkan çok boyutlu yapısal örüntüleri keşfetmek için analitik bir araç olarak ele almaktadır.

## **2. LİTERATÜR TARAMASI**

Yapay zekanın benimsenmesi, iş ve bilgi sistemleri literatüründe ortaya çıkan bir araştırma odağı haline geldi. İlk

başta, araştırmalar teknolojik altyapı ve organizasyonel yeteneğe odaklanmıştı, ancak son birkaç yılda yapay zekanın ve büyük veri yönetiminin benimsenmesi daha karmaşık, çok yönlü bir mesele olarak ortaya çıktı.

Teknoloji–Organizasyon–Çevre (TOÇ) modeli, firmaların yeni teknolojileri benimseme kararlarını açıklamada önemli bir kuramsal zemin sağlamıştır. Model, teknolojik kapasite, organizasyonel hazırlık ve çevresel baskıların etkileşiminin benimsemeyi şekillendirdiğini varsayar (Alsheibani & Cheung, 2018). TOÇ modeline dayanan çalışmalar, YZ benimsemesini etkileyen başlıca faktörlerin veri altyapısı kalitesi, liderlik desteği, kültürel açıklık ve düzenleyici ortam olduğunu göstermektedir (Badghish & Soomro, 2024). Yeniliklerin Yayılımı Teorisi (Diffusion of Innovation – YIT) ise teknolojik yeniliklerin bir organizasyon içinde nasıl yayıldığını açıklamakta ve algılanan fayda, uyumluluk ve karmaşıklık gibi değişkenlerin benimseme hızını etkilediğini öne sürmektedir (Abulail & Badran, 2025). Bu iki kuramın ortak yönü, yapay zekâ benimsemesini yalnızca teknik bir karar değil, aynı zamanda örgütsel davranışın bir yansıması olarak değerlendirmeleridir.

Ampirik literatür, yapay zekâ teknolojilerinin firmaların yenilik, verimlilik ve sürdürülebilirlik performansını güçlendirdiğini göstermektedir. Shen ve Li (2025), Çin’de YZ kullanımının kurumlarda çevresel, sosyal ve yönetim performansını artırdığını bulgulamıştır; Xiong ve Xu (2025) ise “sorumlu yapay zekâ” uygulamalarının uzun vadeli inovasyon kapasitesini güçlendirdiğini belirtmiştir (Xiong & Xu, 2025).

YZ’nin artık belirli sektörlerle özgü değil, “genel amaçlı teknoloji” olarak kabul edildiğini gösteren bulgular da giderek artmaktadır. Bahoo ve Cucculelli (2023), yapay zekânın üretim, finans, sağlık ve hizmet sektörlerinde benzer biçimde benimsenerek rekabet avantajı yarattığını göstermiştir. Benzer

şekilde, Agrawal ve Ahmad (2024), bulut tabanlı altyapıların küçük ve orta büyüklükteki firmaların da büyük firmalarla eşdeğer yapay zekâ kapasitelerine ulaşmasını sağladığını öne sürmüştür.

Ancak, bu literatürdeki çalışmaların çoğu anket veya regresyon analizine dayalıdır ve firmalar arasındaki benimseme örüntülerini veri madenciliği teknikleriyle keşfetmeyi amaçlamamaktadır. Oysa veri madenciliği, özellikle kümeleme analizi (K-Means), firmaları benzer benimseme profillerine göre sınıflandırarak yapay zekâ olgunluğunu nesnel biçimde inceleme fırsatı sunar (Poll & Beijer, 2022). Bu nedenle, bu çalışmada veri madenciliği temelli bir yaklaşım tercih edilmiştir. Bu yaklaşım, yapay zekâ benimsemesinin yalnızca “neden” sorusuna değil, aynı zamanda “nasıl” ve “hangi örüntüler içinde” gerçekleştiği sorusuna da yanıt vermeyi amaçlamaktadır.

### **3. YÖNTEM**

Bu çalışma, firmaların YZ ve özellikle BDM’ni benimseme düzeylerini çok boyutlu biçimde incelemek amacıyla tasarlanmış nicel ve veri madenciliği temelli keşifsel bir araştırmadır. Temel metodolojik yaklaşım, firmaların YZ/BDM kullanım örüntülerini belirlemek için kümeleme analizi (K-Means) yöntemini kullanmak üzerine kuruludur.

#### **3.1. Araştırma Tasarımı**

Araştırma, nedensel bir model kurmaktan ziyade, veri içerisindeki benimseme örüntülerini (adoption patterns) keşfetmeyi hedeflemiştir. Bu yönüyle, çalışma “veri madenciliği” yaklaşımının tanımlayıcı analiz gücünden yararlanarak, firmaların yapay zekâ olgunluğunu ölçülebilir kümeler biçiminde sınıflandırmaktadır (Poll & Beijer, 2022).

Kuramsal arka plan olarak, Teknoloji–Organizasyon–Çevre (TOÇ) modeli ve Yeniliklerin Yayılımı Teorisi (YYT) kavramsal düzeyde dikkate alınmıştır. Bu iki çerçeve, firmaların teknolojik yenilikleri benimseme süreçlerinde; teknolojik kapasite, örgütsel hazırlık ve çevresel baskıların birlikte rol oynadığını vurgulamaktadır (Alsheibani & Cheung, 2018; Badghish & Soomro, 2024). Ancak bu çalışma, bu çerçeveleri istatistiksel test etmeyi değil, veri temelli kümeleme yoluyla ampirik örüntüleri tanımlamayı amaçlamaktadır.

### **3.2. Veri Seti ve Kapsam**

Çalışmada kullanılan veriler, Kaggle platformunda yayımlanan açık erişimli “Global AI Tool Adoption Across Industries and Regions (2023–2025)” veri setinden alınmıştır. Bu veri tabanı, ChatGPT, Claude, Bard, Midjourney ve Stable Diffusion gibi önde gelen üretken yapay zekâ araçlarının ülke, sektör ve firma büyüklüğü bazında benimsenme eğilimlerini temsil etmektedir.

Veri seti sentetik biçimde oluşturulmuş olmakla birlikte, değişkenler arası tutarlılık korunmuş ve analitik olarak “araştırma düzeyinde” bir modelleme düzeyi sağlamaktadır. Benzer sentetik veri tabanlarının, büyük ölçekli dijital dönüşüm ve YZ benimseme araştırmalarında geçerli sonuçlar ürettiği önceki çalışmalarda da doğrulanmıştır (Kim & Lee, 2025).

Veri setinde yer alan temel değişkenler şunlardır:

- Benimseme oranı (%): YZ aracını aktif olarak benimseyen firma oranı,
- Günlük Aktif Kullanıcılar: YZ araçlarını günlük bazda kullanan kullanıcı sayısı,
- Ülke: Firmanın bulunduğu ülke,
- Sektör: Faaliyet gösterilen sektör,

- Firma Büyüklüğü: Firma ölçeği (girişim, KOBİ, kurumsal),
- Yapay Zekâ Aracı: Kullanılan yapay zekâ aracı.

Bu değişkenler, çok boyutlu bir veri yapısı oluşturmak amacıyla sayısal ve kategorik olarak incelendi.

### **3.3. Veri Ön İşleme Süreci**

Analiz öncesinde veri kalitesini artırmak için kapsamlı bir veri temizleme ve standartlaştırma süreci uygulandı. Bu yaklaşım, yapay zekanın benimsenmesini sadece nedensel ilişkiye dayalı olarak değil, aynı zamanda belirlenen davranışsal fenomenlere göre sınıflandırma fırsatı sunar. Bu süreç aşağıdaki adımlardan oluşmuştur:

- Eksik ve uç değerlerin incelenmesi: Eksik gözlemler kaldırılmış; uç değerler Winsorization yöntemiyle düzeltilmiştir.
- Standardizasyon: Benimseme oranı ve günlük aktif kullanıcılar değişkenleri z-skoru yöntemiyle normalize edilmiştir. Bu işlem, ölçek farklılıklarının kümeleme sonuçlarını bozmasını önlemiştir.
- Kategorik değişkenlerin kodlanması: Ülke, sektör, firma büyüklüğü ve yapay zekâ aracı değişkenleri nominal ölçeklerde label encoding yöntemiyle kodlanmıştır.

Bu ön işleme adımları, verinin homojenliğini sağlamış ve K-Means algoritmasının öklidyen uzaklık temelli yapısına uygun hale getirmiştir.

### **3.4. Kümeleme Analizi ve Uygulama Adımları**

Bu araştırmada, firmaların YZ ve BDM kullanım örüntülerini belirlemek amacıyla K-Means kümeleme algoritması uygulanmıştır. K-Means yöntemi, benzer özellikler taşıyan

gözlemleri, belirli sayıdaki kümelere (k) atayarak, her küme içindeki benzerliği en yüksek, kümeler arası farkı ise en düşük düzeyde tutmayı hedefler.

Küme sayısının (k) belirlenmesi için silhouette katsayısı (Silhouette Coefficient) yöntemi kullanılmıştır. Bu yöntem, gözlemler arasındaki içsel tutarlılığı ölçmekte ve kümelerin ne kadar iyi ayrıştığını değerlendirmektedir. Aşağıda, farklı k değerleri için elde edilen silhouette skorları verilmiştir:

**Tablo 1. Küme Sayısına Göre Silhouette Katsayıları**

| Küme Sayısı (k) | Silhouette Katsayıları |
|-----------------|------------------------|
| 2               | 0.3585                 |
| 3               | 0.3742                 |
| 4               | 0.3658                 |
| 5               | 0.3672                 |

En yüksek silhouette değeri  $k = 3$  için elde edilmiştir (0.3742). Dolayısıyla analizlerde üç kümeli çözüm tercih edilmiştir.

Bu üç küme, firmaların yapay zekâ ve büyük dil modeli teknolojilerini benimseme düzeylerine göre sınıflandırılmıştır:

- Yüksek benimseme,
- Orta düzey benimseme,
- Düşük benimseme profilleri.

Kümeleme sonuçları, yalnızca benimseme oranlarına göre değil; aynı zamanda günlük kullanım yoğunluğu, ülke, sektör, firma ölçeği ve kullanılan YZ aracı değişkenleriyle çapraz analiz edilmiştir. Böylece YZ benimsemesi, tek boyutlu bir oran ölçümünden çıkarılarak, çok boyutlu bir olgunluk profili olarak ele alınmıştır.

### **3.5. Analitik Yaklaşımın Gerekçesi**

Veri madenciliği temelli analizlerin, özellikle K-Means kümeleme yönteminin, teknoloji benimseme çalışmalarında

yapısal örüntüleri keşfetmede güçlü bir araç olduğu literatürde vurgulanmaktadır (Poll & Beijer, 2022).

Bu yöntem, yapay zekâ benimsemesini yalnızca nedensel ilişkiler üzerinden değil, gözlenen davranış kalıpları üzerinden sınıflandırma imkânı sunar.

Bu çalışmada K-Means algoritmasının seçilme nedenleri şunlardır:

- Verinin hem sürekli (benimseme oranı, kullanım yoğunluğu) hem kategorik (sektör, büyüklük) bileşenler içermesi,
- Büyük veri yapılarında yüksek hesaplama verimliliği sunması,
- Benimseme olgunluğu gibi gizli (latent) örüntülerin keşfine olanak tanınması.

Bu bağlamda, çalışma klasik istatistiksel testlerin ötesine geçer ve bu bağlamda davranışları sınıflandırmak için desen tanıma kullanarak yapay zekanın benimsenme kalıplarını arar.

## **4. BULGULAR**

### **4.1. YZ ve BDM Benimseme Profilleri**

YZ ve BDM Benimseme Profilleri Küme analizi yoluyla, şirketlerin YZ ve BDM uygulama yoğunluğuna ilişkin üç farklı benimseme profili Tablo 2’de belirlenmiştir. Bu profiller, şirketlerin teknolojiyi benimseme düzeylerini gösterir ve sadece daha geniş ve derin bir benimseme seviyesinde değil, günlük kullanım yoğunluğu açısından da farklılıklar görülmektedir.

**Tablo 2. Yapay Zekâ ve Büyük Dil Modeli Benimseme Profillerine Göre Küme Ortalamaları**

| <b>Benimseme Profili</b> | <b>Benimseme Oranı (%)</b> | <b>Günlük Aktif Kullanıcı Sayıları</b> |
|--------------------------|----------------------------|----------------------------------------|
| Yüksek Benimseme         | 68.46                      | 2,466.36                               |
| Orta Benimseme           | 28.76                      | 4,962.44                               |
| Düşük Benimseme          | 68.05                      | 7,659.51                               |

Tablo 2 incelendiğinde, kümeler arasında adoption rate (benimseme oranı) ve daily active users (günlük aktif kullanıcı) değerleri bakımından anlamlı farklılıklar bulunmaktadır. Yüksek benimseme oranına sahip firmalarda bile günlük kullanım düzeylerinin farklılık göstermesi, firmaların YZ/BDM teknolojilerini stratejik düzeyde benimseyip operasyonel düzeyde henüz tam olarak entegre edememiş olabileceklerini düşündürmektedir.

Bu sonuç, yapay zekâ benimsemesinin yalnızca bir yatırım kararı değil, aynı zamanda operasyonel olgunluk ve organizasyonel dönüşüm gerektiren çok aşamalı bir süreç olduğunu ortaya koymaktadır (Palade & Cârutaşu, 2023). Kümeler genel olarak üç karakteristik profili yansıtmaktadır:

**Düşük Benimseme:** Bu gruptaki firmalar, yapay zekâ teknolojilerini kurum genelinde yaygın olarak benimsememiştir; ancak belirli departmanlarda veya uzman ekiplerde yoğun ve düzenli kullanım gözlenmektedir. Bu yapı, genellikle pilot proje veya deneme aşamasındaki firmaları temsil etmektedir.

**Orta Benimseme:** Bu gruptaki firmalar, YZ/BDM entegrasyonuna stratejik düzeyde yatırım yapmıştır; ancak operasyonel düzeyde kullanım sıklığı ve kapsamı sınırlıdır. Bu durum, benimseme niyetinin yüksek, ancak uygulama olgunluğunun düşük olduğunu göstermektedir.

**Yüksek Benimseme:** Bu gruptaki firmalar, YZ teknolojilerini sistematik bir şekilde işlerine entegre etmiş ve bunları günlük faaliyetlerde etkili bir şekilde uygulamışlardır.

Analiz sonuçları, benimseme oranı ile kullanım yoğunluğunun her zaman paralel ilerlemediğini göstermektedir. Bu bulgu, firmaların yapay zekâ teknolojilerini benimseme sürecinde stratejik karar alma ile operasyonel uygulama arasında belirli bir zaman farkı bulunduğunu ortaya koymaktadır. Literatürde “YZ olgunluğu (AI maturity)” olarak tanımlanan bu durum, yapay zekâ adaptasyonunun yalnızca teknolojik bir yatırım değil, aynı zamanda insan sermayesi, süreç uyumu ve örgütsel öğrenme gerektiren bir dönüşüm süreci olduğunu desteklemektedir (Neumann & Guirguis, 2022).

#### **4.2. Ülkeler Arası Karşılaştırma**

YZ ve BDM teknolojilerinin benimsenmesinde ulusal düzeydeki yüksek homojenlik, analiz sonuçlarından iyi bir şekilde anlaşılmaktadır. Analiz edilen tüm ülkelerde, kabaca, şirketlerin üçte biri yüksek benimseme profiline sahipken, benzer bir oranı orta benimseme profiline sahiptir. Buna karşılık, düşük benimseme grubundaki şirketlerin oranı ülkeler arasında biraz farklılık göstermektedir.

**Tablo 3. Ülkelere Göre Yapay Zekâ ve Büyük Dil Modeli Benimseme Profillerinin Dağılımı**

| Ülke                        | Yüksek Benimseme | Orta Düzey Benimseme | Düşük Benimseme |
|-----------------------------|------------------|----------------------|-----------------|
| Avustralya                  | 0,358            | 0,365                | 0,278           |
| Brezilya                    | 0,364            | 0,355                | 0,281           |
| Kanada                      | 0,354            | 0,365                | 0,280           |
| Çin                         | 0,352            | 0,363                | 0,285           |
| Fransa                      | 0,354            | 0,362                | 0,284           |
| Almanya                     | 0,361            | 0,361                | 0,278           |
| Hindistan                   | 0,361            | 0,360                | 0,279           |
| Güney Kore                  | 0,360            | 0,364                | 0,276           |
| Birleşik Krallık            | 0,370            | 0,357                | 0,273           |
| Amerika Birleşik Devletleri | 0,360            | 0,359                | 0,281           |

Tablo 3 incelendiğinde, ülkeler arasındaki sapmaların çok çok küçük olduğunu göstermektedir. Dolayısıyla, yüksek

benimseme oranı %35–%37 civarında değişirken, düşük benimseme oranı %27–%28 aralığında kalmaktadır. Bu bulgu, yapay zekâ ve büyük dil modellerinin benimsenmesinin yalnızca gelişmiş ekonomilere özgü bir fenomen olmadığını; aksine, bulut tabanlı teknolojiler ve erişilebilir dijital altyapılar sayesinde farklı ekonomik bağlamlarda faaliyet gösteren firmalar arasında küresel ölçekte yaygınlaştığını göstermektedir.

Özellikle Birleşik Krallık (%37,0) ve Brezilya (%36,4) gibi farklı ekonomik yapıya sahip ülkelerde yüksek benimseme oranlarının benzer düzeylerde görülmesi, YZ ve BDM teknolojilerinin “genel amaçlı teknoloji” niteliğini desteklemektedir. Ayrıca, Çin, Hindistan ve Güney Kore gibi Asya ekonomilerinde gözlemlenen dengeli dağılım, bu ülkelerdeki ulusal yapay zekâ stratejilerinin firmalar düzeyinde somut yansımalar yarattığını göstermektedir.

Sonuç olarak, ülke bazlı analizler, YZ/BDM benimsemesinin artık yalnızca teknolojik olgunluk göstergesi değil, aynı zamanda küresel dijitalleşme sürecinin ortak bir paydası hâline geldiğini ortaya koymaktadır. Bu durum, firmaların coğrafi konumdan ziyade örgütsel strateji, dijital altyapı ve insan sermayesi yeterlilikleri doğrultusunda farklılaşmaya başladığını göstermektedir.

#### **4.3. Sektörel Dağılım**

Sektör bazlı analizler, YZ ve BDM benimsenme profillerinin sektörler arasında büyük ölçüde dengeli bir dağılım gösterdiğini ortaya koymuştur. Elde edilen bulgular, YZ ve BDM teknolojilerinin yalnızca bilgi yoğun sektörlerde değil, aynı zamanda geleneksel üretim alanlarında da yaygınlaştığını göstermektedir.

**Tablo 4. Sektörlere Göre Yapay Zekâ ve Büyük Dil Modeli Benimseme Profillerinin Dağılımı**

| <b>Sektör</b> | <b>Yüksek Benimseme</b> | <b>Orta Benimseme</b> | <b>Düşük Benimseme</b> |
|---------------|-------------------------|-----------------------|------------------------|
| Tarım         | 0.359                   | 0.364                 | 0.277                  |
| Eğitim        | 0.364                   | 0.357                 | 0.279                  |
| Finans        | 0.356                   | 0.363                 | 0.281                  |
| Sağlık        | 0.362                   | 0.363                 | 0.275                  |
| Üretim        | 0.356                   | 0.361                 | 0.283                  |
| Perakende     | 0.357                   | 0.358                 | 0.285                  |
| Teknoloji     | 0.357                   | 0.364                 | 0.278                  |
| Ulaşım        | 0.364                   | 0.359                 | 0.277                  |

Tablo 4 incelendiğinde, sektörler arasındaki farkların oldukça sınırlı olduğu görülmektedir. Yüksek benimseme oranları %35–36 bandında değişmekte, düşük benimseme oranları ise %27–28 düzeyinde seyretmektedir. Bu sonuç, YZ ve BDM teknolojilerinin yalnızca belirli uzmanlık veya bilişim odaklı sektörlerde değil, yatay biçimde tüm sektörlerde benzer ölçüde benimsenen “genel amaçlı teknolojiler” hâline geldiğini göstermektedir.

Özellikle teknoloji (%35,7), finans (%35,6) ve sağlık (%36,2) sektörlerinde beklenen yüksek benimseme oranlarının, tarım (%35,9), üretim (%35,6) ve ulaştırma (%36,4) gibi geleneksel sektörlerle benzer seviyelerde seyretmesi dikkat çekicidir. Bu durum, YZ/BDM'nin sektör bazlı bir avantajdan kurumsal dönüşümün ortak bir boyutuna geçtiğini göstermektedir. Eğitim ve sağlık gibi kamu hizmeti sektörlerinin yüksek benimseme oranı, YZ'nin sadece bir organizasyon veya firmanın yararına değil, aynı zamanda verimliliği, erişimi ve hizmet kalitesini artırmak için bir yenilik sağlayıcısı olarak görüldüğünü de göstermektedir. Dolayısıyla, AI ve BDM teknolojilerinin tüm sektörlerde yaygın olarak benimsenmesi, teknoloji transferi aşamasından geçtiklerini ve sektör yapısından bağımsız olarak kurumsal dönüşüm için evrensel bir araç haline geldiklerini ortaya koymaktadır. YZ ve BDM'nin benimseme

profilleri, şirket büyüklüğünden bağımsız olarak oldukça benzerdir. Şirket büyüklüğüne göre yapılan analizler, AI ve BDM benimseme profillerinin geniş bir şekilde dağıldığını önermektedir. Gerçekten de şirket büyüklüğü ile AI/BDM benimseme seviyesi arasında bir ayrımın olmaması, bu teknolojilerin demokratikleştiğini, erişilebilir ve ölçek bağımsız olduğunu ortaya koymaktadır.

#### **4.4. Firma Büyüklüğüne Göre Benimseme Profilleri**

Firma büyüklüğü bazında yapılan analizler, YZ ve BDM benimseme profillerinin firma büyüklüğünden bağımsız olarak oldukça benzer bir dağılım gösterdiğini ortaya koymuştur. Bu bulgu, teknolojinin yalnızca büyük ölçekli kurumsal yapılarda değil, aynı zamanda küçük ve orta ölçekli işletmeler (KOBİ) ile girişim ekosisteminde de benzer düzeylerde benimsendiğini göstermektedir.

**Tablo 5. Firma Büyüklüğüne Göre Yapay Zekâ ve Büyük Dil Modeli Benimseme Profillerinin Dağılımı**

| <b>Firma Büyüklüğü</b> | <b>Yüksek Benimseme</b> | <b>Orta Benimseme</b> | <b>Düşük Benimseme</b> |
|------------------------|-------------------------|-----------------------|------------------------|
| Kurumsal               | 0.362                   | 0.359                 | 0.279                  |
| KOBİ                   | 0.360                   | 0.360                 | 0.280                  |
| Girişim                | 0.356                   | 0.365                 | 0.279                  |

Tablo 5 incelendiğinde, üç farklı ölçek grubunda da yüksek benimseme oranlarının %35–36 aralığında seyrettiği, orta ve düşük düzey benimseme oranlarının da birbirine oldukça yakın olduğu görülmektedir.

Bu sonuç, YZ ve BDM teknolojilerinin ölçeklenebilir ve bulut tabanlı yapılarının, teknolojiye erişimdeki geleneksel “ölçek avantajı” farklarını önemli ölçüde ortadan kaldırdığını göstermektedir. Artık yüksek bilişim altyapısı yatırımı gerektirmeyen SaaS (Software as a Service) ve API tabanlı YZ çözümleri sayesinde, küçük ölçekli firmalar da büyük firmalarla benzer teknolojik kapasiteye erişebilmektedir.

Girişim firmaları açısından görece daha yüksek “orta düzey benimseme” oranı (%36,5), bu firmaların yenilikçi ancak sınırlı kaynaklı bir teknoloji benimseme stratejisi izlediğini göstermektedir. Bu profil, genellikle deneme ve uyarlama (trial–adaptation) aşamasındaki teknolojik girişimleri yansıtmaktadır.

Büyük ölçekli firmalarda ise yüksek benimseme oranının (%36,2) yanında düşük kullanım oranlarının varlığı, stratejik düzeyde yatırım kararlarının alınmış olmasına rağmen operasyonel yayılım sürecinin zamana yayıldığını göstermektedir.

Bu sonuç, yapay zekanın artık büyük sermaye firmalarının tekelinde olmadığını, aksine her türden şirketin rekabet gücünü artırmak için bir araç olduğunu göstermektedir. YZ araçlarına göre benimseme YZ araçlarına dayanan bir çalışma, BDM ve üretken YZ platformlarına dayalı AI araçları arasında benimseme seviyelerinde net bir fark bulamamaktadır.

#### **4.5. Kullanılan YZ Araçlarına Göre Benimseme**

YZ aracı bazlı analizler, farklı BDM ve üretken YZ platformları arasında benimseme düzeyleri açısından anlamlı bir farklılaşma olmadığını göstermektedir. Elde edilen bulgular, firmaların belirli bir YZ aracını tercih etme kararının markaya veya platforma olan bağlılıktan ziyade, kurumsal strateji, kullanım amacı ve entegrasyon düzeyine dayandığını ortaya koymaktadır.

**Tablo 6. Kullanılan Yapay Zekâ Araçlarına Göre Benimseme Profillerinin Dağılımı**

| <b>BDM Aracı</b> | <b>Yüksek Benimseme</b> | <b>Orta Benimseme</b> | <b>Düşük Benimseme</b> |
|------------------|-------------------------|-----------------------|------------------------|
| Bard             | 0.356                   | 0.366                 | 0.278                  |
| ChatGPT          | 0.363                   | 0.359                 | 0.278                  |
| Claude           | 0.353                   | 0.372                 | 0.274                  |
| Midjourney       | 0.359                   | 0.361                 | 0.280                  |
| Stable Diffusion | 0.354                   | 0.361                 | 0.285                  |

Tablo 6 incelendiğinde, tüm araçlarda yüksek benimseme oranlarının %35–36 bandında, orta düzey benimseme oranlarının ise %36–37 aralığında seyrettiği görülmektedir. Bu sonuç, araçlar arasındaki farkların istatistiksel olarak önemsiz düzeyde olduğunu göstermektedir.

Özellikle ChatGPT (%36,3) ve Claude (%37,2) kullanıcılarının görece yüksek benimseme oranları, bu araçların kurumsal süreçlere entegrasyon kolaylığı ve çok yönlü kullanım olanaklarıyla ilişkilendirilebilir. Ancak Bard, Midjourney ve Stable Diffusion gibi araçların yayılması, tüm YZ araçlarının kurumsal ekosistemlerde eşit derecede değerli görüldüğünü ortaya koymaktadır. Bu sonuç, YZ ve BDM teknolojilerinin marka merkezli gelişmelerden, kurumsal iş akışına dahil edilen sıradan üretkenlik araçlarına evrildiğini önermektedir. Ayrıca, bu sonuç, YZ benimsemesiyle ilgili belirleyici faktörün artık kullanılan araç değil, teknolojinin organizasyon yapısı ve süreçleriyle entegre edilip edilmediği olduğunu göstermektedir. Araç bazlı analizler, bu nedenle, kurumsal YZ olgunluğunun teknik tercihten çok yönetim stratejisi ve dijital dönüşüm kapasitesi ile daha fazla ilişkili olduğunu ortaya koymakta, YZ/BDM teknolojilerinin "teknoloji markaları" yerine kurumsal yetkinlik göstergeleri olarak algılanması gerektiğini daha da kanıtlamaktadır. Tartışma YZ ve BDM teknolojisi bağlamında firmaların genel benimseme profilleri, bu çalışmada ülke düzeyinde ve farklı sektörler ve şirket büyüklükleri arasında analiz edilmiştir. Sonuçlar genel olarak, YZ/BDM teknolojilerinin küresel ölçekte dostça, ölçeklenebilir ve sektörler arası yatay bir dönüşüm aracı haline geldiğini göstermektedir.

## **5. TARTIŞMA**

Bu çalışmada, firmaların YZ ve BDM teknolojilerini benimseme profilleri, ülkeler, sektörler ve firma ölçekleri

düzeyinde kapsamlı biçimde incelenmiştir. Bulgular genel olarak, YZ/BDM teknolojilerinin küresel ölçekte erişilebilir, ölçeklenebilir ve sektörler arası yatay bir dönüşüm aracına dönüştüğünü göstermektedir. Bu bölümde elde edilen bulgular, mevcut literatürle ilişkilendirilerek tartışılmıştır.

Araştırma bulguları, firmalar arasında yüksek benimseme oranlarına rağmen günlük kullanım düzeylerinde önemli farklılıklar olduğunu göstermektedir. Bu durum, yapay zekâ benimsemesinin yalnızca bir stratejik karar değil, operasyonel olgunluk gerektiren çok aşamalı bir dönüşüm süreci olduğunu ortaya koymaktadır. Literatürde de benzer biçimde, YZ'ye hazırlık (AI-readiness) ve YZ olgunluğu kavramları kurumsal dönüşüm sürecinin kademeli doğasını vurgulamaktadır. Kurumların YZ teknolojilerini yalnızca satın almakla kalmayıp, iş süreçlerine entegre etmesi gerektiği; bunun ise veri yönetimi, çalışan yetkinliği ve süreç adaptasyonu gibi boyutlara bağlı olduğu belirtilmektedir (Palade & Căruțașu, 2023). Benzer şekilde, Neumann & Guirguis (2022) yapay zekâ benimsemesinde teknolojik faktörlerin ilk aşamalarda belirleyici olduğunu; ancak organizasyonel kültür, üst yönetim desteği ve çalışan tutumlarının süreç ilerledikçe daha kritik hâle geldiğini ortaya koymuştur. Bu bağlamda çalışmamızın bulguları, benimseme olgunluğu kavramını destekler niteliktedir.

Ülke bazlı analizler, YZ/BDM benimseme oranlarının gelişmiş ve gelişmekte olan ekonomiler arasında büyük farklılıklar göstermediğini ortaya koymuştur. Bu bulgu, yapay zekânın küresel erişilebilirliğinin arttığını ve “teknolojik eşitsizlik” kavramının zayıfladığını göstermektedir. Bu sonuç, literatürdeki TOÇ çerçevesine dayalı yaklaşımlarla da uyumludur. TOÇ modeli, çevresel faktörlerin (örneğin regülasyonlar, piyasa baskısı, teknolojiye erişim) firmalar arası benimseme düzeylerini etkilediğini, ancak bulut tabanlı altyapıların bu farklılıkları giderek azalttığını göstermektedir

(Alsheibani et al., 2018). Ayrıca, Badghish & Soomro (2024) ayrıca, küçük ve orta ölçekli işletmelerin (KOBİ'ler) bile çevresel ve insan sermayesi yetkinlikleri sonucunda YZ teknolojisini benimsemeye ileri düzey firmalara yaklaştığını eklemektedir. Bu çerçevede, çalışmamızın ülkeler arası yakınsama bulgusu, dijital altyapıların ve bulut tabanlı yapay zekâ araçlarının küreselleşmenin yeni itici gücü hâline geldiğini desteklemektedir.

Sektör bazlı analizler, YZ/BDM teknolojilerinin yalnızca bilgi yoğun alanlarda değil, tarım, üretim ve ulaştırma gibi geleneksel sektörlerde de benzer benimseme düzeylerine ulaştığını göstermektedir. Bu bulgu, yapay zekânın “General-Purpose Technology (GPT)” belirli bir alana özgü olmayan, çok çeşitli endüstrilerde verimliliği artırabilen bir dönüşüm aracı olduğunu doğrulamaktadır. Benzer biçimde, Khanfar & Mavi (2024) yaptıkları sistematik analizde, YZ'nin endüstriyel sınırları aşan, tüm sektörlerde verimlilik ve karar destek kapasitesini güçlendiren yatay bir teknolojiye dönüştüğünü belirtmiştir. Ayrıca, Alzahrani (2025) sektörler arasında farklılıkların daha çok teknolojik olgunluk ve veri altyapısı düzeylerinden kaynaklandığını, ancak genel olarak yapay zekânın “sektör bağımsız rekabet avantajı” sunduğunu vurgulamaktadır.

Firma büyüklüğüne göre yapılan analizlerde, büyük ölçekli firmalar, KOBİ'ler ve girişim firmaları arasında benimseme profillerinin şaşırtıcı derecede benzer olduğu görülmüştür. Bu durum, teknolojinin demokratikleşmesi kavramıyla açıklanabilir. Agrawal & Ahmad (2024)'e göre, küçük firmalar, bulut altyapısı ve SaaS tabanlı AI çözümleri tarafından yönlendirilen ölçek avantajı boşluklarını hızla kapatmakta, böylece firma tarafında rekabet eşitliğini artırmaktadır. Bu bakış açısından, çalışmamızın sonuçları, yapay zekanın yalnızca büyük sermaye firmalarının alanından, ölçekten bağımsız bir üretkenlik ve yenilik sürücüsüne dönüştüğünü göstermektedir. YZ araç bazlı analizlerin karşılaştırılması,

organizasyonların hedefleri, kullanım durumları ve entegrasyon stratejilerinin, teknoloji seçimlerinde marka kadar önemli olduğunu, ChatGPT, Claude, Bard ve diğerleri arasındaki yakın benimseme oranlarının da bu durumu desteklediğini göstermektedir. Bu bulgu, Albanese (2024)'in bulgularına benzemektedir; yazara göre, şirketler YZ araçlarını seçerken markaya olan bağlılıktan çok organizasyonel hedeflere uyum, veri güvenliği ve entegrasyon kolaylığına odaklanmaktadır. Dolayısıyla, YZ / BDM teknolojilerinde, kurumsal stratejik uyum ve dijital dönüşüm yeteneği, marka farklılığından daha önemlidir.

## **6. SONUÇ VE ÖNERİLER**

Bu çalışma, firmaların YZ ve BDM teknolojilerini benimseme düzeylerini veri madenciliği temelli bir yaklaşımla analiz etmiş; ülkeler, sektörler, firma büyüklükleri ve kullanılan yapay zekâ araçları düzeyinde çok boyutlu bir değerlendirme sunmuştur. Araştırmanın temel amacı, firmalar arasında YZ/BDM benimseme örüntülerini keşfetmek ve bu örüntülerin kurumsal olgunluk, stratejik yönelim ve dijital altyapı kapasitesiyle ilişkisini ortaya koymaktır.

Bulgular, küresel ölçekte oldukça homojen bir benimseme yapısı bulunduğunu göstermiştir. Ülkeler arasında yüksek benimseme oranlarının %35–37 bandında seyretmesi, YZ teknolojilerinin artık sadece gelişmiş ekonomilere özgü bir fenomen olmaktan çıkıp dijital altyapının erişilebilirliği sayesinde küresel düzeyde demokratikleştiğini göstermektedir. Benzer biçimde, sektörler ve firma ölçekleri arasında gözlemlenen minimal farklar, YZ/BDM'nin “genel amaçlı teknoloji” niteliğini desteklemekte; teknolojinin, üretimden finansa, eğitimden sağlığa kadar çok çeşitli alanlarda yatay bir dönüşüm aracı hâline geldiğini göstermektedir.

Araştırma, aynı zamanda benimseme oranı ile kullanım yoğunluğunun her zaman paralel ilerlemediğini ortaya koymuştur. Bu durum, firmaların YZ teknolojilerini stratejik düzeyde benimsese de operasyonel entegrasyonun zaman içinde olgunlaştığını göstermektedir. Başka bir ifadeyle, YZ benimsemesi bir “teknoloji satın alma süreci” değil, bir “organizasyonel öğrenme ve kültürel dönüşüm sürecidir. Bu sonuç, "yapay zekâ olgunluğu" literatüründe savunulan kademeli olgunlaşma yaklaşımı ile paraleldir.

### **6.1. Araştırmanın Akademik Katkıları**

Bu çalışma literatüre üç açıdan katkı sağlamaktadır:

**Veri madenciliği temelli bir model kullanması:** Çoğu önceki araştırma anket veya regresyona dayalıyken, mevcut çalışma, gözlemlenen verilere dayalı olarak firma davranışını sınıflandırmak için bir K-Means kümeleme yöntemi kullanmaktadır.

**Benimseme örüntülerini çok boyutlu ele alması:** YZ/BDM benimsemesi yalnızca oranlarla değil, aynı zamanda günlük kullanım yoğunluğu, sektör, ülke ve büyüklük etkenleriyle birlikte değerlendirilmiştir.

**TOÇ ve YYT kuramlarını ampirik olarak yeniden yorumlaması:** Bu modellerin öngördüğü teknolojik, örgütsel ve çevresel faktörlerin veri temelli kümelenmelerde nasıl biçimlendiği ortaya konmuştur.

### **6.2. Uygulayıcılar ve Politika Yapıcılar İçin Öneriler**

Elde edilen bulgular, politika yapıcılar ve işletme yöneticileri için de önemli çıkarımlar sunmaktadır:

Dijital altyapı yatırımları, artık yalnızca büyük ekonomilerde değil, gelişmekte olan pazarlarda da YZ benimsemesini hızlandırıcı etki yaratmaktadır. Bu nedenle kamu politikaları, bulut tabanlı sistemlere erişimi kolaylaştıran

düzenlemelere ve açık veri paylaşım mekanizmalarına öncelik vermelidir.

Firma yöneticileri açısından YZ/BDM uygulamalarında başarı, yalnızca yazılım seçimiyle değil, insan sermayesi yatırımı, veri yönetimi standartları ve değişim yönetimi süreçleriyle doğrudan ilişkilidir.

KOBİ ve girişim firmalar, büyüklük avantajı eksikliğini SaaS tabanlı yapay zekâ çözümleriyle aşabilmektedir. Dolayısıyla, girişimcilik destek programlarının YZ tabanlı üretken araçlara erişimi teşvik etmesi, rekabet eşitliğini güçlendirecektir.

### **6.3. Gelecek Araştırmalar İçin Öneriler**

Bu çalışma keşifsel nitelikte olduğundan, elde edilen bulguların nedensel ilişkilerle test edilmesi ileriki çalışmalar için önemli bir adım olacaktır. Önerilen çalışma, gelecekteki çalışma için aşağıdaki önerileri içermektedir:

- Gerçek (sentetik olmayan) mikro veri setlerinin kullanımı, sonuçların dış geçerliliğini artıracaktır.
- Boylamsal veri analizleri, şirketlerde zaman içinde YZ olgunluğu hakkında içgörü elde etmek için harika bir yaklaşım olacaktır.
- Teknolojik, organizasyonel ve çevresel faktörlerin göreceli etkilerinin sayısal testleri, nedensel modelleme (örneğin, yapısal eşitlik modelleme veya PLS-SEM) temelinde yapılabilir.
- Kültürel, etik ve düzenleyici boyutların YZ benimsemesi üzerindeki etkisi, özellikle farklı bölgelerdeki sosyo-kültürel faktörlerin nasıl rol oynadığı açısından incelenmeye değerdir.

- Son olarak, büyük dil modellerinin örgütsel bilgi yönetimi üzerindeki mikro etkileri (ör. çalışan verimliliği, karar kalitesi, yaratıcı üretim süreçleri) niteliksel yöntemlerle derinleştirilebilir.

Genel olarak, bu çalışma yapay zekâ ve büyük dil modellerinin artık sadece teknolojik bir inovasyon değil, kurumsal dijital olgunluğun belirleyici unsuru hâline geldiğini ortaya koymuştur. YZ/BDM teknolojileri, firma büyüklükleri, sektör veya coğrafyadan bağımsız biçimde evrensel bir dönüşüm altyapısı oluşturmakta; dolayısıyla firmaların rekabet avantajını şekillendiren yeni bir paradigma yaratmaktadır. Bu sonuç, yapay zekânın gelecekte yalnızca üretkenlik ve verimlilik artışı değil, aynı zamanda örgütsel öğrenme, etik yönetim ve kapsayıcı dijital dönüşüm süreçlerinin merkezinde yer alacağını göstermektedir.

## **KAYNAKÇA**

- Abulail, M., & Badran, A. (2025). Revisiting the diffusion of innovation theory in the context of artificial intelligence adoption: Perceived benefit and complexity as determinants. *Journal of Technology and Innovation Management*, 42(1), 44–61.
- Agrawal, R., & Ahmad, S. (2024). Cloud-based AI adoption and the democratization of digital transformation among SMEs. *International Journal of Information Systems and Business Transformation*, 15(3), 210–227.
- Albanese, D. (2024). Enterprise-level selection of generative AI tools: Strategic alignment and integration priorities. *AI and Business Review*, 9(2), 88–102.
- Alzahrani, F. (2025). Artificial intelligence adoption across industries: Determinants of sectoral variation in digital maturity. *Technology in Society*, 75, 102489.
- Alsheibani, S., & Cheung, Y. (2018). A framework for AI adoption in organizations: Integrating the TOE model. *Procedia Computer Science*, 138, 534–541.
- Badghish, S., & Soomro, T. (2024). Artificial intelligence adoption by SMEs to achieve digital transformation: A TOE perspective. *Information Systems Frontiers*, 26(2), 455–472.
- Bahoo, S., & Cucculelli, M. (2023). Artificial intelligence and corporate innovation: A cross-industry perspective. *Technological Forecasting and Social Change*, 194, 122563.
- Khanfar, A., & Mavi, R. (2024). Artificial intelligence as a general-purpose technology: A systematic review of industrial applications. *Journal of Industrial Engineering and Management*, 17(1), 23–47.

- Kim, J., & Lee, D. (2025). Synthetic data and model validity in large-scale AI adoption research. *Data Science and Society*, 6(1), 15–33.
- Kou, G., & Tang, Y. (2025). Clustering analysis of enterprise AI adoption patterns: Insights from data-driven modeling. *Expert Systems with Applications*, 242, 122957.
- Neumann, T., & Guirguis, M. (2022). Adoption of artificial intelligence in organizations: The role of culture, leadership, and employee attitudes. *Journal of Organizational Change Management*, 35(6), 1189–1208.
- Palade, V., & Căruțașu, G. (2023). The role of artificial intelligence in organizational digital maturity and governance. *Management & Marketing: Challenges for the Knowledge Society*, 18(2), 250–265.
- Poll, H., & Beijer, J. (2022). Clustering methods for digital transformation maturity: Applications in AI adoption analysis. *Information and Management Research Journal*, 61(4), 301–319.
- Sánchez-Calderón, C., López-Fernández, M., & Gutiérrez, A. (2025). Integrating TOE and DOI models to explain AI adoption in SMEs. *Computers in Human Behavior Reports*, 15, 101276.
- Shen, W., & Li, F. (2025). Artificial intelligence adoption and corporate ESG performance: Evidence from Chinese firms. *Journal of Cleaner Production*, 445, 140982.
- Xiong, J., & Xu, Y. (2025). Responsible AI practices and firm innovation performance: Evidence from global enterprises. *Technovation*, 134, 102893.

# YÖNETİM BİLİŞİM SİSTEMLERİ ALANINDA BİLİMSEL ARAŞTIRMALAR

**yaz**  
yayınları

YAZ Yayınları  
M.İhtisas OSB Mah. 4A Cad. No:3/3  
İscehisar / AFYONKARAHİSAR  
Tel : (0 531) 880 92 99  
yazyayinlari@gmail.com • www.yazyayinlari.com