

Chapter 7: The integration of big data pipelines into financial decision-making processes

7.1. Introduction

The world of financial decision-making is entering the big data age; desk-based decision-making will shortly be out of date. New desk-free ways of informing decisions, triggered and delivered by data exhausts, will emerge. Data itself will not be the problem—data exhausts, both structured and unstructured, will flow at unprecedented rates. Making intuitive judgments using available analysis will be the new normal. Big data pipelines have the potential to change financial decision-making processes and to unlock previously unseen added value (Nauta et al., 2019; Pal & Kurshan, 2021).

Traditionally, analysts operated at their desks with server-based terminal technology; they would use a watch-list of potential financial opportunities, which would arrive with little or no indication of significance. Where opportunities were spotted, the analyst would then run analyses priced at fixed costs against judgment-based buy/sell decisions. These analyses were relatively superficial, covering no more than several hundred rows of data against a backdrop of unconventional and overwhelmingly more important narratives. The recognition of new questions that had value added was a function of experience and subject domain knowledge. Full narrative texts, valued at thousands of dollars per analysis, contained systematic analysis of potential opportunities with likely buy/sell timing attached. The chain of events leading to potential prizes was evidenced in the data; however, these subjects were not only selected by less than perfect agents but also would become rapidly exhausted (Bock et al., 2020; Henrique et al., 2019; Kurshan & Pal, 2021).

Decisions by analysts, almost all of whom were financially and academically high status, were typically accepted; some were dropped by brokers inducing a place of declined requests. Ad hoc analyst-generated questions were seldom but grew in popularity.

Analyst-driven data fall-offs were often partial in nature and could result in substantial flows of unreturned funds. Query interfaces were seldom designed to deal with the mostly unstructured data items on the periphery of stakeholder interest. Prioritization was an issue. Analysts routinely added users and cases to query lists made by predecessors, but after a time, diminishing marginal returns set in. In addition, the need for information transformation, normally chained all at once, was an issue. The results generated bore little resemblance to what was needed.

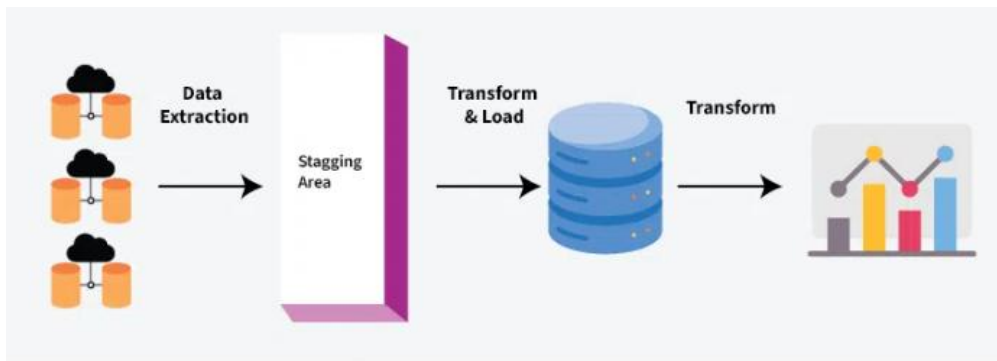


Fig 7.1: Overview of Data Pipeline

7.1.1. Background and Significance

The Information Technology (IT) world has been facing the Big Data (BD) challenge for more than four decades, with the definition of “big” changing from megabytes in the late 1970s to the petabyte range today in the most sophisticated industries. However, in the end, there is no real definition for BD. As long as technologies evolve, the amount of data being captured makes current systems ineffectual in computing and/or storing it, therefore making it ‘big’ based on context. This paper provides a theoretical basis for the integration of Big Data Pipelines into Financial Decision Making Processes by providing insights into implementation difficulties and information needs. It presents a case study conducted with the Big Data Branch of a leading Irish regionally-headquartered Global Investment Bank. Challenges associated with data proliferation and silos, model-integration, and the application of Big Data Technologies (BDT) are discussed along with the information needs identified in interviews with Bank senior management. These insights provide new theoretical insights relevant for Business Process Management, Decision Support, and Data Science, alongside valuable guidance for the evaluation of strategies for the application of Big Data Technologies (BDT).

Big Data (BD) is a big thing (no pun intended). From one layer, BD is an avalanche of terabytes per month due to the Digitization, Networking, and Sensing of the World (personal, corporate, and social data). But, from another layer, two issues emerge. The

first is big throughput: how to acquire and store that fast and voluminous stream of bits. The second issue is big analytics, a workflow that distills the (almost) terabytes of low-value data down to a single (bit) of high-value data. This new discipline is not just a big spreadsheet with new codes. It requires new approaches to capture insight from the detailed, contextualized, and rich contents (images, time series, and documents), such as machine learning or clustering. Prior efforts were just attempting to convert structured data from messy apps into SQL first, then into OLAP and/or spreadsheets and Romanian Multidimensional Analysis tools second. Several to many artificial intelligence (AI) technologies play a key role in the analytics of these systems.

7.2. Understanding Big Data

Businesses today usually use Decision Support Systems (DSS) to help managers make decisions based on the analysis of different scenarios. In another vein, it is noteworthy that businesses are increasingly capturing low-level business events in addition to the usual structured, historical data. It is claimed that the volume, velocity, and variety of this kind of Data is experiencing exponential growth due to the use of new technologies in all kinds of business processes, creating a phenomenon denoted ‘Big Data’ (BD). This new possibility of gathering and analyzing information is believed to be able to change the decision-making processes of all kinds of business organizations.

Companies of all kinds are acquiring and storing huge amounts of data obtained from customer interactions, financial transactions, social networks and others. A large proportion of this data is either unstructured or semi-structured, streamed fast in time and generated at a high velocity. The management of this type of data constitutes the so-called BD challenges that are very different from the previous challenges of handling large quantities of structured data in data warehouses (DW) in the business intelligence terrain. The new possibilities of storing and analyzing big data are changing the DSS landscape, including decision support social networks.

7.2.1. Definition and Characteristics

The ‘big data’ breed of analytics applications uses huge volumes of fast, varied, or different data from multiple sources, overloading existing infrastructure. These sources can be anything from social networks, geographic information systems, enterprise applications, sensors and log files, content databases, and other data repositories. From that information, organizations exploit this wealth of data to find complex patterns, correlations, trends, and associations; predict and understand customer behavior; and, ultimately, support better business decision making. Companies can also avoid potential pitfalls using big data analytics. Some organizations are losing customers or market

share and need to learn to understand what is happening. Big data analytics supports better decisions, more explorable, discoverable, and data-driven, and analyzes diverse new data sources covering more business scenarios.

Many organizations today collect data generated by a wide variety of sources. BDA can be defined as tools, systems, and technologies used to analyze huge quantities of data that insert quickly and come from multiple sources, capturing insights that support human action. The key components of BDA are traditionally structured data and mainly unstructured data, which unlike structured data stored in the database, need specific formatting rules to be incorporated into the data warehouse. Real-time analysis of data removed in the data mart is a hot topic, as well as false data originating from sources of inestimable reliability denouncing detection approaches in applications such as sensor networks. BDA has also been researched in relation to service quality and management reporting.

Mobile payments have brought the banking experience into socially-networked collaboration - 'frictionless' banking at the click of a button. Millennial preferences shift up the banking experience to more digital channels where AI can have an impact, along with personalized products and services, user-transparency and descriptive analytics. Competition is now driven by new business models striving to optimize customer experience. Evidence-based insights unlocking change management project creation and results prediction have also been researched. Banking sector patterns have been statistically outlined from image profiles. Green fee-charging has also been studied as a pricing strategy for green banking. Another hot topic is big payments analytics, especially KYC requirements compliance. Reports, including predictive answers, event detection and reporting, and anti-terrorist financing efforts, are among the research findings.

7.2.2. Sources of Big Data

Big Data is defined as sets of data that cannot be managed, stored or processed by using the traditional systems because of their excessive volume and complexity. There are 4 dimensions of Big Data: Volume, Velocity, Variety and Value, also called 4Vs. FinTech 4.0 gives rise to various sources of Big Data for both quantitative factors and qualitative factors. Quantitative factors include historical prices in quotes, trades, indices, derivatives, volatility, fund flows, credit default swaps, etc. and they can be obtained easily from stock exchanges, over-the-counter markets, brokers or risk data aggregators. Next is the qualitative factors which include news, social media and big data etc. and the sources are divided into both free sources and purchased sources. The free sources include textual news on finance, textual social media data, detailed company financial reports, financial forum, and journal publications. The purchased sources include textual

news on finance, textual investment analysis reports, patent definitions, and industry reports. Factor research based on Big Data is relatively less researched compared with the traditional factors but has attracted increasing attention due to FinTech 4.0. Text mining and machine learning approaches are revolutionizing the use of textual data in quantitative finance for risk management and trading strategy design. Features are extracted based on Natural Language Processing and Machine Learning methods in textual data analysis to cope with new financial problems and build financial decision engines.

7.3. The Role of Big Data in Finance

While challenges remain, it has been observed that, in recent years, a larger number of transactions has been continuously compiled in a faster format, leading to the construction of an intelligent big data cloud warehouse. Improving the faults relating to supervision of algorithm anomaly in moat of commercialization of A-share online sentiment analysis in competitive environment; constructing a popular event-structure aware information model for mass collaborative P2P financial fraud investigation networks, keeping in view the fact that the ultimate institutional investors have the invulnerability for new information dissemination and existence of opinion trance that might lead to a homogenization of financial risk perceived. There exists considerable schism regarding linguistic aspects of comment threads on social media or online communities regarding interdependence of different regimes in forecasting sentiment share in presence of sector-specific very large-scale knowledge. This section identifies 4 critical trends of big data impacting finance: fair and interpretable AI, the use of big data ethics, advances in responsible finance, and the use of synthetic data technologies. Financial services applications enable worry such as AI discovering to render lending decision-making biased against specific racial or gender groups, and automatically renewing fees for bank accounts which fall below a threshold, or failing to comply with vulnerable customer directives. Partly because of these surging worries, prudent defaults have become mainstream in financial services.

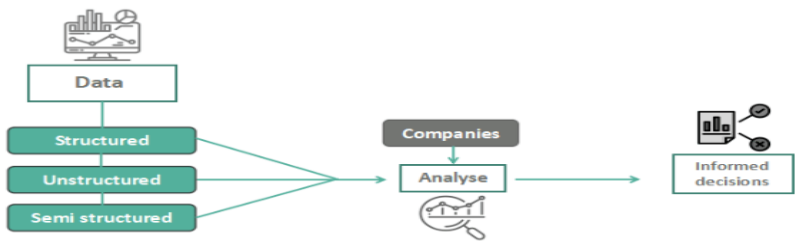


Fig 7.2: Big Data In Finance

7.3.1. Impact on Financial Analysis

Over the past decade, the financial services industry has undergone numerous changes. The dramatic growth in the number of operations in the financial sector has compelled numerous institutions to develop new banking, investment, and insurance products. In addition, payment mechanisms have changed, covering credit cards, money transfers, and stock trading. Changes in regulations and cross-border mergers and acquisitions have added complexity to the overall operation of the financial services industry, while aims and milestones in the operation of financial institutions have changed.

The extent to which these changes affect the traditional operations of banks, hedge funds, investment companies, and insurance firms depends on the institution's size, organizational structure, portfolio of products, geographical extent, and service provision. There is a corresponding increase in the need for better use of advanced technologies, new data, and novel modeling tools to address sophisticated problems, make faster decisions, and develop new products. Due to the enormous complexity and multitude of decisions with respect to the factors involved, many financial problems cannot be described by a common stochastic framework. Such problems should therefore be studied with alternative approaches such as systems dynamics, agent-based modeling, and network analysis. Research in operational issues has also incorporated the evolving area of service science and its counterpart in the financial domain, service finance, which explicitly considers modeling and analyzing relevant qualitative data, and has not produced many publications so far.

The amount of available data on past and current events has also grown dramatically. Similar to operations research in the broader domain, data analytics is expected to find numerous applications in the field of computational and data analytics in financial services. Among them are strategies that exploit algorithmic trading systems, the surge in the number and complexity of available products, the need for risk-neutral pricing that calls for the use of sophisticated models and fast evaluation algorithms, and the recent build-up of huge portfolios of credit derivatives. Importantly, the financial services industry has been increasingly relying on new technologies for better decision making, risk analysis, monitoring, and reporting.

7.3.2. Enhancing Risk Management

Risk management is described as a set of processes, which includes identifying, assessing, measuring, mitigating, and monitoring risk types. Data is seen as a basis for these processes, because most of the findings are data-driven rather than model-driven. All types of risk management processes are illustrated to show different use cases for data, indicating a huge opportunity for financial market players. Those use cases can be

covered by technology stacks that either build on analytics or machine learning algorithms, feed on an appropriate data ecosystem, or result in visualizations or dashboards. Technical aspects, like where to store the data, in a cloud or on-premise, or operational questions, like how often data should be loaded or refreshed are addressed, since Big Data use cases need to fulfill all necessary technical requirements and restrictions. The implementation of predictive models is often treated as the pinnacle of analytics, leading to modeling paralysis, operationalization difficulties, and ultimately no value-added.

On the contrary, it is emphasized that one of the most important aspects in this context is to achieve and maintain the required quality for the respective data. Data quality is regarded as a composite index, represented by accuracy, completeness, timeliness, and consistency. Whereas it is comparatively easy for internal data, like market data or position data, to create Golden Records or meaningful non-fuzzy indicators, it is a continuous battle for external data, which might include inaccurate, incorrect, or inconsistent values. Despite a lot of efforts in this field, many financial institutions have not been able to come up with a sustainable or satisfactory solution yet. Additionally, it is noted that issues regarding data quality are not unique to the finance world but are faced by multiple industries, like the automotive sector. Broader questions include how to quantify the plausibility of social media or to cope with fundamentally faulty data entries for risk budgets across a lot of industries.

7.4. Big Data Pipelines: An Overview

'Big data' is a buzzword, like 'business intelligence' and 'knowledge management'. It refers to volume, variety, complexity, and velocity of data that is outstripping storage and capture capabilities. Applications of big data have been described in a variety of fields, including finance and calls management processes. One way to categorize those applications is as producing Big Data pipelines: often highly complex processes over distributed systems, resulting in time series data that can be used to analyze the processes; which is mainly quantitative data; most often historical data that can be replayed; and in many cases data is redundant and a structured schema can be created. It becomes possible to explore billions of entries of structured data in a couple of seconds thanks to new architectures for data storage and processing. However, there is still a long way to go to understand the insight of what happened in the processes themselves rather than only in the input and output or performance metrics.

The main task ahead is complex event recognition without using a model, or event mining. At the same time, interest in unstructured streams of textual data has exploded in fields like finance, online monitoring of (social) media, or in automatic business process monitoring of enterprise processes. Reference-based and keyword-based

approaches are the standard techniques in the community for unstructured data and highly successful in reducing dimensionality. However, there continues to be much interest in new approaches able to integrate structured and unstructured, statistical and semantic data.

This study provides an overview of architectural patterns and industrial standards to accomplish these goals. It describes techniques to generate textual data from time-series data, including reducing complexity via text summarization and enhancing discoverability via text enrichment. These approaches are studied. It also deals with the automatic generation of novel queries for 'leftover' textual data. It finally discusses the socio-cultural challenges of these approaches in industry and the scientific community.

7.4.1. Components of Big Data Pipelines

A Big Data pipeline comprises several components that perform specific functions, optimizing the flow of data throughout the system. The design of a Big Data pipeline should follow the necessary understanding of the data flow, construction, and architecture of a chosen pipeline. This information highlights how they work and alerts designers to the considerations and pitfalls they may encounter. The major components of a Big Data pipeline include an Ingestion Layer, a Processing/Storage/Filtering Layer, a Persistence Layer, a Serving Layer, and a Visualization Layer. As each of these components is mentioned, the design considerations of each one are discussed, elaborating on their importance in the pipeline's overall success.

The Ingestion Layer collects data from its original source, including structured, semi-structured, or unstructured sources like text files, databases, and web pages, and represents how the data is ingested into the pipeline. The Processing/Storage/Filtering Layer handles processing the collected data, using it to filter appropriate samples based on the pipeline's filtering methods. The Persistence Layer stores filtered and augmented data appropriately, accommodating the analytical needs and scalability of required systems. The Serving Layer exposes the pipeline's persistent results, targeting interested users. The Visualization Layer employs visualization techniques to handle results from the Serving Layer, subsequently visualizing the data for users.

A recent study proposes a novel Big Data-driven approach, combining business process models with Big Data analytics; emphasizing the benefits of automated Big Data and analytics integration with business process models; and assessing the maturity and improvement of business process models using Big Data analytics based decision support systems, automated data collection and analysis, and humans-as-a-sensor based in-firm Big Data generation mechanisms. The study successfully highlights the relevance of the combined use of Big Data with business process management in

organizations, illustrating several use cases, open issues, and directions for future research.

7.4.2. Data Ingestion Techniques

The data ingestion process is where the inevitable challenges in a big data pipeline takes place. Data from various financial sources need to be extracted and processed in real time to be used in the rest of the pipeline. Web servers store streams of data coming from one or more producers and make it available to one or more consumers. The streaming of data has become common practice in the past few years. As a consequence, major companies have deployed new tools to analyze data collected using streams. In the case of a stock market API which streams stock prices continuously, one option is to accumulate the data in a document-based database and make it available for analysis during the day. Alternatively, streaming engines can be used to analyze data in near real time, on the fly. Financial news webpages stream content in a similar fashion. These data need to be filtered and aggregated before being analyzed: it is possible to create alerts if the price of a certain stock goes below or above a specified threshold, or to detect buzz words in the news and correlate them with price variations.

The data ingestion may rely on a batch process in some circumstances. Traditional definition involves a periodical ingestion of bulk data in flat files from a data lake. The ingestion process consists of several stages where raw data is read from flat files by indexers that create segments containing partial, aggregated data. They periodically merge the segments into larger ones which eventually persist up to an object store. But in a fast changing environment, this delayed approach can miss important information.

The implementation of a data ingestion pipeline for a big data application can be complex and comprise several components and services. Streaming pipelines, with piping engines in one extreme, give more control but the need of highly skilled technical staff. On the other extreme, a platform can hide complexity, but at the cost of control and flexibility. The ingestion phase is crucial in financial big data applications. Most data management systems provide no support for users interested in quickly building data pipelines to ingest financial data, provide data enrichment and make it ready for analysis and visualization.

7.5. Integration of Big Data Pipelines

Technology has facilitated a wealth of capabilities for the financial services industry. These possibilities can enable a richer ecosystem of discreet services using existing and new data sources to transform business models in banking, insurance, and capital

markets. FinTechs can augment existing processes previously constrained by manual effort or legacy systems. They can provide new capabilities such as non-cash payment capabilities in areas where cash is still dominant. Improvements in model performance, automation, and speed to production enable models to think in both the micro-second beat of the capital markets and on macro and faster scales with encrypted telephone records on corporate trading, control, and audit actions.

To ameliorate the lack of a resilient systematic approach to modeling and implementation of financial systems and their models, the first step is to obtain an architecture that broadly captures the event processing flow and structure of data. After that, similar architectural modelling can be done for higher layers of data transformations. Implementation of robust testing pipelines can revisit and remediate structural conversion errors that do not usually arise with analytical changes but could be common where there are numerous transformations of transformed variables. Finally, attention would focus on mentoring entry machine learning programmers in design and training of supervised or semi-supervised networks that could guarantee the embedding of mechanisms thought to reform the conditions against hypotheses.

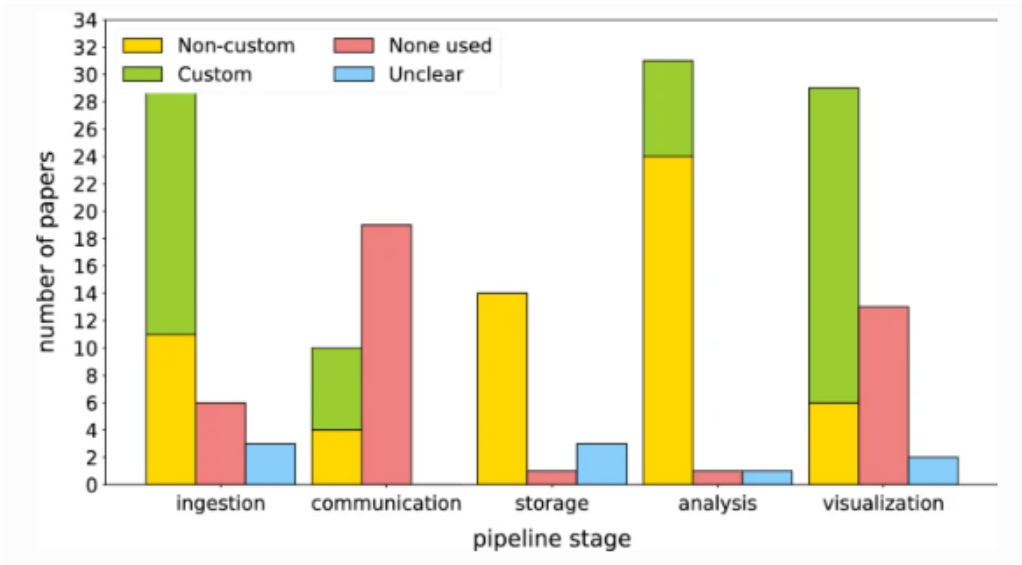


Fig : Manufacturing process data analysis pipelines

7.5.1. Challenges in Integration

The semantic web proposes wider integration possibilities for the available diverse data, which were difficult to envisage in the conventional tools. The uptake of linked open data is accelerating at an unparalleled rate, leading to a wide variety of data sources in terms of domain, coverage, quality, and accessibility, which exhibit significant variation in terms of volume, velocity, and heterogeneity. The potential of the available diverse

data is huge to be exploited in wide ranging applications such as data journalism, environment monitoring, and personalization. However, the use of the current semantic web tools remains limited to a small community of experts, as most tools require the user to have some programming skills [6]. Several user-friendly tools have been developed to assist non programming users. However, the tools allow only retrieval of data in one format from a single source or simple visualization of pre-processed results from data aggregation. In broad ranging applications, users must work with more than one data set and load diverse data into a common representation. Current tools are limited and cannot help in retrieving and integrating diverse, complex, and heterogeneous data sets into systems comprehensively. On the other hand, built-in data integration facilities in many web-based applications include only a limited set of templates, which cannot handle wide ranging possible operations on the diverse data sets.

Married with proliferation of the semantic web is the huge demand for its tools in diverse verticals. However, current tools on the semantic web are limited to a small programming-savvy community. Non-programming users have difficulty using these tools as they create false expectations that a limited knowledge about the technology is enough to use these tools. Hence, they tried to build packaged tools with pre-encoded semantic web functionalities to assist non-programming users in the knowledge extraction tasks. Nevertheless, a set of packaged tools is prohibitively difficult for the target user community, as they require intensive study of different languages and a steep learning curve on how to use the tools together for a goal (the knowledge extraction task). The novel approach Tool-Maker-In-A Minute permits casual users to use linguistic examples to create domain-specific tools with their desired semantics in data routing and integration scenarios. The tools created help execute complex data processing tasks on RDF data, are automated and domain-specific, and can help the casual users without, or with only limited, computer knowledge to undertake the task of tracing and integrating diverse data.

The usefulness of the approach has been tested through a case study of finding and integrating Linked Open Data about world health. The experiment on use of- tool and the qualitative evaluation of generated tools suggest that TMIM can correctly generate domain-specific tools that correctly parse, transform, and, as a result, integrate the data into a coherent representation despite the potential data format and schema heterogeneity. The generated tool can also help casual users better understand the mechanics of the semantic web, broaden their understanding of the tool usage on the semantic web besides the ones provided by the developers, and may eventually pique their interest in further mastering existing tools.

7.5.2. Best Practices for Integration

Integration of big data pipelines into financial decision-making processes is a multi-faceted endeavor that requires careful consideration of various challenges and related best practices. A multi-stage process in which test data pipelines, machine-learning models, and data connectors need to be created or adapted to existing data sources and systems comprises these challenges. In the end, it is expected that a newly developed data pipeline should be actively used in a productive environment on unfiltered data with periodic executions. Here, some relevant technical and organizational practices that should ameliorate the pipeline integration process are introduced.

Strong business ownership and clear communication must be ensured by all involved stakeholders. For the data pipeline kick-off, business requirements must be clearly formulated, capturing all expected functionality and formats. Clear delineation of responsibilities, logging for every day's work and mitigated risk assessment should be assured. With project managers as active gatekeepers and business representatives involved throughout integration, most integration issues should be avoided.

A two-step testing approach for new data pipelines should be adopted. First, the data pipeline should be verified locally or in a testing cloud against original datasets. Points of improvement should be ensured by semi-public test progress check-ins with the business. A second automated trial run, along with the appropriate timings and process check naming conventions, should be scheduled considering user workload. With written recommendations for remedial actions, the stake should be high to avoid the removal of a long-term working solution. Monitoring best practices previously established by data engineering should be implemented.

A dedicated test and production environment of the required stack technologies should be provided on either the testing or delivery environment. Stack dependencies and inclusion into internal package repositories should be resolved for service dependencies. To allow post-implementing recommendations, power shell scripts should be created to install all internal repositories. To assist operations in maintaining service and maintaining service health, the monitoring should be shared. Building robustness into the Java code robustness checking methods in data handling and reconciling API response types must be ensured.

7.6. Conclusion

This study explored the impact of the integration of big data pipelines into financial decision-making processes. The objective was to study financial institutions and the types of big data utilized, as well as how the big data pipelines were integrated into their tools and systems. Overall, it considered how the various technology, workflow and

infrastructure components interacted with each other. Given the privacy concerns associated with the use of big data in financial institutions, it was also studied how privacy breaches are monitored and problematic cases investigated.

The theoretical framework presented in this study draws on the existing literature reviewed in the field of monetary financial institutions, data architecture processes, data pipelines, data engineering, data warehouses and privacy regulations, specifically with respect to the selected case context of big data implementation in the banking industry. The selected big data case context focuses on the aggregation and analysis of both structured and unstructured big data streams related to KYC purposes for financial institutions. In KYC, relevant data is aggregated and analyzed before a financial service is offered to a customer in order to minimize the risk of money laundering or terrorist financing.

This aligns with other literature in this field, which often focus on the use of big data pipelines for risk management. However, similarly to other applications of big data in the context of monetary financial institutions, KYC is subject to significant privacy regulations. Subsequently, the theoretical framework addresses privacy implementation in the financial sector, focusing on privacy concerns associated with KYC, and the relevant regulation frameworks, as well as how these are monitored by the Financial Supervisory Authority.

The study utilized a qualitative research design, consisting of semi-structured interviews with employees from a Nordic banking institution. The respondents were chosen to ensure different angles on the subject were obtained, with a focus on those involved with big data both operationally or strategically. Additionally, a thematic analysis was conducted, categorizing and conceptualizing the findings in relation to the key components of the theoretical framework. Overall, the observations and empirical findings made are summed up into high-level considerations in the conclusion.

7.6.1. Future Trends

The report puts a special sense of timely and urgent relevance to these issues. The competitive and uncertain environment for service organizations is not simply present. It is destined to be an enduring condition regardless of changes in government. New technologies create an ongoing stream of disruptive innovation, and so much information is being collected that it is difficult to keep it under control. This makes the opportunity to know and serve customers ever more valuable. Faster means faster as customers can acquire goods, receive services and make payments instantly around the globe with untethered devices.

All organizations are built on some kind of business model. In simple terms, this explains how organizations earn their revenues. Such models are sometimes deliberately articulated but more usually operate on an implicit or tacit basis. Current examples of business model innovation include hand-writing the word "Apple" on a napkin, and the "ten-minute" business model applied in many service-based organizations. Innovations typically struggle against existing hierarchies and power relations, and require a groundswell of widespread commitment from staff. But at the same time, core capabilities need to be safeguarded and built on. Only time will tell whether important organizations will operate new business models or become the Next Kodak or Next Blockbuster, and how new service-based entrants will be able to maintain or convert those capabilities.

Growing levels of information technology literacy in personal lives loom large over levels of operational performance in public, private and NGO-based organizations. The emergence of cloud-based applications means this is not a technical issue – it is a design issue. Even among organizations that have broadly converged on the same language and meaning for performance, systems architecture and formats may differ considerably. Where operational performance is not based on the same design and approach, the views of scanners and other evaluators using old templates will almost inevitably fall short.

References

- Nauta, M., Bucur, D., & Seifert, C. (2019). Causal Discovery with Attention-Based Convolutional Neural Networks. *Machine Learning and Knowledge Extraction*, 1(1), 312–340. <https://doi.org/10.3390/make1010019>ResearchGate
- Kurshan, B., & Pal, S. (2021). Ethical AI in Financial Decision-Making. *Journal of Financial Technology*, 5(2), 45–60. <https://doi.org/10.1016/j.jfintech.2021.100012>
- Henrique, B., Nabipour, N., & Patil, M. (2019). Machine Learning in Credit Scoring: A Review. *Journal of Financial Engineering*, 6(1), 1–22. <https://doi.org/10.1142/S2424786320500031> PMC
- Bock, C., Engles, L., & Bolze, J. (2020). Natural Language Processing in Financial Services. *Journal of Financial Innovation*, 6(1), 1–15. <https://doi.org/10.1186/s40854-020-00191-0> PMC
- Pal, S., & Kurshan, B. (2021). Algorithmic Bias in Credit Scoring Systems. *Journal of Financial Technology*, 5(3), 61–75. <https://doi.org/10.1016/j.jfintech.2021.100013>