

POLITECNICO DI TORINO

Collegio di Ingegneria Gestionale

**Master of Science Course in
Decision making for technology and social change**

Master of Science Thesis

Analysis of the Applications of Artificial Intelligence in the Supply Chain



**Politecnico
di Torino**

Tutor

Prof. CARLO RAFELE

Prof. MASSIMO REBUGLIO

Candidate

Matteo Chirichilli

OCTOBER 2025

Contents

INTRODUCTION	5
CHAPTER 1 – What is Artificial Intelligence.....	6
1.1 Definition of AI	6
1.1.1 The Roots of the Definition.....	6
1.1.2 Artificial Intelligence Today: An Open Definition	6
1.2 History of AI	7
1.2.1 The Birth of AI (1950s).....	7
1.2.2 The Years of Expectations and Limits (1960s–1970s)	7
1.2.3 The Revival through Machine Learning (1980s–1990s).....	7
1.2.4 The Era of Deep Learning (2010s–Present).....	8
1.3 Key Technologies Today	9
1.3.0 Introduction.....	9
1.3.1 Machine Learning	9
1.3.2 Deep Learning.....	9
1.3.3 Natural Language Processing (NLP)	10
1.3.4 Computer Vision	11
1.3.5 Autonomous Robotics.....	11
1.4 Evolution of AI in the Industrial Context	11
1.4.1 Early Applications of AI in Industry (1980s–1990s)	11
1.4.2 Introduction of Machine Learning and Big Data (2000s–2010s).....	11
1.4.3 The birth of Smart Factories (2020s–Present).....	12
CHAPTER 2 – Supply Chain: Structure and Challenges	13
2.1 Introduction to Supply Chain	13
2.1.1 Definition	13
2.1.2 Historical Evolution: From Traditional Logistics to the Modern Supply Chain	13
2.2 Supply Chain structure: stakeholders and main processes	14
2.2.1 Main Stakeholder.....	14
2.2.2 Main processes.....	14
2.3 System for Supply Chain Management.....	15
2.3.1 Physical flow systems	15
2.3.2 Information flow systems.....	16
2.3.3 Financial flow systems.....	17
2.4 Current challenges.....	17
2.4.1 Demand volatility.....	17
2.4.2 Environmental and social sustainability.....	18
2.4.3 Risk management and resilience	19
2.5 Summary and outlook	20
CHAPTER 3 – AI’s applications in Supply Chain	21
3.1 Introduction	21
3.1.1 Chapter goals	21
3.1.2 Types of Artificial Intelligence in the Supply Chain	21

3.2 Demand Forecasting.....	22
3.2.1 Definition.....	22
3.2.2 Predictive Models vs. Traditional Approaches.....	22
3.2.3 Machine Learning models.....	25
3.2.4 Deep Learning models	27
3.3 Inventory optimization with AI	29
3.3.1 Definition and Challenges.....	29
3.3.2 Traditional Inventory Management Models.....	30
3.3.3 Metaheuristic Optimization in Inventory Management	31
3.3.4 Deep Reinforcement Learning	33
3.4 Logistics	34
3.4.1 Introduction.....	34
3.4.2 Vehicle Routing Problem (VRP).....	35
3.5 Computer Vision for Quality Control.....	38
3.5.1 Introduction.....	38
3.5.2 Machine Vision System's components.....	39
3.5.3 AI Techniques for Quality Control.....	40
3.5.4 Colour Recognition and Measurement.....	40
3.5.5 The importance of AI in Quality Control	41
3.6 Summary and Outlook	41
CHAPTER 4 – AI Applications in Supply Chains.....	43
4.1 Introduction.....	43
4.2 Walmart sales forecasting.....	43
4.2.1 Case study introduction.....	43
4.2.2 Data exploration.....	43
4.2.3 Modeling methodology	44
4.2.4 Models results	45
4.2.5 Managerial implications, limitations and future prospects	46
4.3 Reinforcement Learning for inventory management.....	47
4.3.1 Introduction.....	47
4.3.2 Reinforcement Learning for Multi-Product and Multi-Node Inventory Management in Supply Chains	47
4.3.4 From Research to Practice: Real-World Evidence	52
4.4 AI-based AGV Applications in Inventory Management: The JD.com Case Study	53
4.4.1 Introduction to JD.com and the competitive environment	53
4.4.2 Technological architecture: AGVs and AI algorithms for inventory management.....	54
4.4.3 Implementation challenges and managerial solutions	56
4.4.4 Future prospects	57
4.5 AI Applications in Last-Mile Delivery: The Canada Post Case Study	57
4.5.1 Introduction and Problem Context	57
4.5.2 System Architecture and Data Preparation.....	58
4.5.3 Deep Learning Models.....	59

4.5.4 Results and Analysis	61
4.5.5 Implications and Conclusions	63
4.6 Practical Application of AI-based Computer Vision in Quality Control	64
4.6.1 Case study introduction and goals.....	64
4.6.2. System Architecture and Experimental Setup	65
4.6.3 Model Training Methodology	66
4.6.4 Results and Performance Evaluation	66
4.6.5 Implications for the Supply Chain and Conclusions	70
CHAPTER 5 - From Case Studies to Strategic Insights: The Future of AI in Supply Chain Management.....	72
5.1 Integrating Lessons from Real-World Applications	72
5.2 - Strategic and Managerial Implications.....	72
5.3 - Challenges and Barriers to AI Adoption for SMEs	73
5.4 - Roadmap for AI Implementation in Supply Chains	74
5.5 - Future Outlook and Research Directions	75
ACKNOWLEDGMENT	77
BIBLIOGRAPHY AND WEBSITE	78

INTRODUCTION

The purpose of this thesis is to identify and study the uses of Artificial Intelligence (AI) in the supply chain, in order to gain a better understanding of how this modern technology has impacted the industrial context. Specifically, it will be explained how the concept of "artificial intelligence" was born and what meaning modern society attaches to this technology; at the same time, a brief history of its use in the supply chain context will be drawn up, highlighting which systems were used during the technological development between the 1950s and today. The basic layout of the supply chain will be outlined next, followed by the technologies used and the difficulties that now jeopardize its efficiency. Following that, an examination of the various AIs employed in the supply chain sectors will be conducted, followed by practical examples of how this technology can be successfully applied in industrial contexts.

CHAPTER 1 – What is Artificial Intelligence

1.1 Definition of AI

1.1.1 The Roots of the Definition

In the last few decades, artificial intelligence (AI) has gone from being a theoretical idea to a real technology that can have a big effect on everyday life, the economy, and the generation of knowledge. Before looking at its uses, effects, or future possibilities, though, it is important to answer a very important question: “What is artificial intelligence?”

John McCarthy, Marvin Minsky, Nathaniel Rochester, and Claude Shannon came up with the term "artificial intelligence" in 1956 during the famed Dartmouth Summer Research Project¹. The first idea was both grandiose and vague: looking into the prospect of robots acting in ways that could be seen as intelligent. This method has led to a very vast field that has changed over time based on different theoretical ideas, new technologies, and research goals.

In 1995, Stuart Russell and Peter Norvig published their groundbreaking book, "Artificial Intelligence: A Modern Approach"². These two authors have worked hard to make the phrase "Artificial Intelligence" clearer. They have changed their definition many times, up to the book's fourth edition in 2020. For a complete understanding of this concept, it is beneficial to utilize the classification suggested by these computer experts, which defines AI along two axes:

- What it does (think or do)
- What it looks like (human vs. rational)

From this taxonomy, four further definitions of Artificial Intelligence arise³:

The first definition talks about systems that can think like people. This group contains systems that can learn and solve issues like a person. For example, models that mimic human memory or analogical reasoning are examples of this type of system. The interest here is not just in the end result, but also in how it was reached.

A second description talks about systems that act like people. Here, the focus is on whether the machine acts like a person, not if it "thinks like a person." The Turing Test is one example. It says that a machine is smart if it can trick a person into thinking it is smart like a human during a conversation.

The third definition is about systems that think logically. A system is smart if it uses formal logic or optimum rationality to make decisions. These kinds of systems can come to the right conclusions based on the right premises, which makes them act like completely rational cognition.

The last description is for systems that act rationally, which means they make the best choices to reach their goals based on the information they have.

Today, most computer scientists are working on making systems that fit this last definition: "Rational Agents" that can see their surroundings and the outside world and then do things that will help them reach the goal for which they were made.

1.1.2 Artificial Intelligence Today: An Open Definition

Today, artificial intelligence is an enabling technology. It works at scale across many fields, from medicine to finance, robotics, and data analysis. Still, its rise hasn't led to a clear or fixed definition. In fact, the more AI blends into daily life, the harder it is to pin down what it actually

is. This lack of clarity isn't a flaw. It reflects the changing nature of the field itself. AI is always shifting. Each new step forward changes what we call "intelligent."

This view helps explain the so-called "AI effect." The term appears often in academic writing. It points to a trend: once an AI-related tool becomes useful, common, and trusted, people stop seeing it as AI. It gets treated as regular or even outdated technology⁴.

In short, artificial intelligence shows our need to build tools that resemble us, help us, or go beyond us. It's not an answer but it's a question. And with each new breakthrough, we ask it again. This thesis will trace the shape of that question by looking at how AI is defined, applied, and governed today, with a focus on the Supply Chain.

1.2 History of AI

1.2.1 The Birth of AI (1950s)

Long before computers, people were working on artificial intelligence. Aristotle in the 4th century BCE inquired if logic might be used to formalize human thought. But these notions didn't really catch on until the 20th century. Alan Turing's 1950 work "Computing Machinery and Intelligence"⁵, was an important turning point. He asked a brave question in it: "Can machines think?" He also came up with the Turing Test as a means to see how smart a machine is.

The 1956 Dartmouth Conference, which was directed by John McCarthy, Marvin Minsky, Nathaniel Rochester, and Claude Shannon⁶, is commonly seen as the beginning of the field. At that time, the term "artificial intelligence" was first used, and the goal was to make machines that could think like people. In the years that followed, early programs were made to answer logic issues or prove arithmetic theorems. One example is Logic Theorist (1956) by Allen Newell and Herbert Simon⁷.

1.2.2 The Years of Expectations and Limits (1960s–1970s)

The AI of the 1960s and 1970s was very excited by symbolic models that were built on clear rules and formal logic. The main notion was that logical structures could be used to show how people think and know things, and that a computer system could work with these structures. In this situation, "expert systems" were created by the scientists. These are systems that try to make the same decisions as an expert person in a small number of specific situations: for example, figuring out what is wrong with a patient or setting up an industrial plant. These systems used a knowledge base (a set of "if-then" rules) and an inference engine that could use those rules to look at data and come to conclusions. MYCIN⁸ is one of the most well-known computer programs. It was built at Stanford University to help doctors figure out if someone has a bacterial infection. Even though these systems did well in certain areas, they had huge problems when it came to handling new, unclear, or incomplete scenarios. This caused a crisis of trust that led to the first "AI winter"⁹.

1.2.3 The Revival through Machine Learning (1980s–1990s)

Since the 1980s, artificial intelligence has gone through a time of new ideas and methods. After being disappointed by the symbolic method, scientists have progressively started to look at ways that can learn from facts. The machine learning paradigm has emerged in this context, poised to transform the development of AI fundamentally. Machine Learning¹⁰ is a type of artificial intelligence that makes algorithms that can detect patterns in data and get better at their jobs over time without having to be trained for each task. Machine learning enables

systems to "learn" directly from experience, adapting to novel or unforeseen information. This is not the same as symbolic AI, which was based on rules set by experts.

At the time the most used algorithms were:

- Support vector machines (SVM) were the most frequent method at the time. Vapnik and Cortes came up with them in the 1990s. They work well for regression and classification in hard-to-understand domains¹¹;
- Decision trees are easy to understand and can help you sort things into groups.
- K-nearest neighbours (k-NN) is a method for finding patterns that doesn't utilize parameters.
- Bayesian methods for modelling probability.

Artificial intelligence is starting to employ the scientific method more as machine learning becomes more popular. This means looking at things, making models, and checking them. It doesn't rely as much on symbolic reasoning. This turning point is an important step that sets the foundation for the significant advances in AI that are happening now, starting with the advent of deep learning in the next ten years.

1.2.4 The Era of Deep Learning (2010s–Present)

The early 2010s were a turning point in the growth of AI, to the extent where they are now seen as the start of the contemporary era of AI. This change happened because three important things came together at the right time:

1. The exponential growth of computational capacity, made possible in particular by the use of GPUs (Graphics Processing Units), originally meant for graphics processing in video games but which have proven to be highly successful in the parallel training of neural networks.
2. The presence of a large quantity of data from diverse sources (for example social networks, sensors, e-commerce, etc.) that is essential for supplying models that require substantial information consumption.
3. Algorithmic innovation, especially the rediscovery and growth of deep neural networks, which are the core of what is known as deep learning.

To comprehend the importance of this pivotal moment, it must be examined the implications of the latest and most important AI development, Deep Learning. It is a type of machine learning that uses artificial neural networks with several layers (This is where the use of the word "deep" comes from). Each layer can get more abstract representations from the data that is fed into it. A network that has been trained to detect photos, for instance, has layers that find lines and edges, other layers that scan for forms and textures, and other one for full things like faces or animals.

The best thing about deep learning is that it can learn important features from raw data on its own, without the need for a person to do it. This method has been very useful for hard and high-dimensional issues, like speech recognition, computer vision, and natural language processing¹².

Since the first steps, deep learning-based models have been used in a lot of different fields:

- Computer vision (detecting objects, recognizing faces, and making diagnostic images).
- Machine translation (Google Translate based on neural models).
- Speech synthesis and recognition (voice assistants like Alexa or Siri).

- Generation and understanding of natural language through large linguistic models (LLM), like GPT, BERT, and their offshoots.

In the years that followed, generative language models have gotten more and more attention, leading to tools like ChatGPT. These systems can write logical documents, answer complex questions, perform creative tasks, and simulate real conversations. One of the most advanced, and contested, areas of modern AI is how they work: they rely on billions of parameters trained on vast collections of text.

This approach is not well seen because no one fully understands how these models produce certain outputs, and their training data often includes content that is biased, misleading, or based on fake news. As a result, the systems can generate answers that sound plausible but if checked they result to be wrong or based on stereotypes. Since users can't trace the source or reasoning behind a response, it becomes difficult to assess its accuracy or intent. These concerns raise questions not just about performance, but about trust, responsibility, and how such tools should be used.

1.3 Key Technologies Today

1.3.0 Introduction

Nowadays Artificial intelligence is a discipline that is articulated through numerous technologies, each of which contributes to the development of artificial cognitive capabilities. The main technologies that today constitute the heart of modern AI are machine learning, deep learning, natural language processing, computer vision and autonomous robotics.

1.3.1 Machine Learning

Machine learning is the most crucial technology that has revolutionized the way we think about AI since the 1980s. When computer scientists develop an algorithm in traditional programming, they write down every step the system takes. Machine learning algorithms instead learn from labelled or raw data and develop their own models to make their predictions.

There are three main ways to teach how to act to the system:

- Supervised learning: the algorithm learns from data sets given by computer scientists with already labels on them. The goal is to find a function that find the proper output when given known inputs. Sorting emails into spam and not spam, or recognizing pictures, is a popular example.
- Unsupervised learning: the data isn't labelled, so the machine has to find patterns or structures on its own. Marketing teams often use function this to group clients by how they act, which is an example of clustering.
- Reinforcement learning: the systems learn by interacting with the environment and gaining rewards or punishments (decided by their developer) for what they do. It helps with things like smart gaming systems and robots that can do things on their own¹³.

1.3.2 Deep Learning

Deep learning is now seen as one of the most important parts of modern AI. It is not only a mechanism to learn with machines but it's a new paradigm that lets artificial systems get, store, and use information in a very hierarchical and scalable way.

Deep artificial neural networks, which are made up of several layers of artificial neurons, are what deep learning is all about. Each layer gives the output created to the next layer that receive it as an input, processes it, and then improves the representation of the data: this is done layer

by layer, some times in a way that is not understandable even for programmers of the system. This multilevel approach lets the system go from simple traits to more complicated ideas, which lets it find deep structures even in input that have outliers, are noisy or not are not structured¹⁴.

There are different deep learning network topologies based on the kind of data and the goal:

- Convolutional Neural Networks (CNNs): These are great for analysing images because they use unique layers to find spatial patterns like edges, textures, and forms. They are what makes computer vision programs work.
- Recurrent Neural Networks (RNNs): made to handle sequences of data, such texts or time series. They have an internal memory that lets them handle the relationship between consecutive pieces of the sequence.

Until now it was highlighted only the winning side of using deep learning, but as every technologies at the beginning of its lifecycle, it has several problems. The most important problems are:

- High data needs: To work well, deep models need a lot of labelled data.
- High computational cost: it takes a lot of energy to train a new model based on deep learning. For instance the energy used to train OpenAI's gpt-4 model could have powered 50 American homes for a century. According to The Economist, to train the next generation of AI Chatbot it will take more or less 1 billion dollar¹⁵.
- Functioning hardly understandable: Deep learning models are like "black boxes," which makes it hard to understand how a decision is reached. This problem is not negligible in fields like justice or healthcare.
- Vulnerability to adversarial attacks: Even very strong models can be fooled by small, imperceptible changes in the input data, which raises problems about security and reliability.

Even if these limitations still exist, deep learning is still one of the most powerful tools available to research and industry today, pushing the boundaries of automation, artificial perception and symbolic intelligence. Its ability to generalize from complex data and learn abstract representations makes it a central technology for all future evolutions of artificial intelligence¹⁶.

1.3.3 Natural Language Processing (NLP)

Natural Language Processing (NLP) is a field of AI that studies how computers and human language interact. Thanks to this technology computer can perceive, interpret, change, and create voice and text. NLP has had a complicated history. It started with systems that used grammatical and syntactic principles, and now it uses neural models that learn directly from language data.

Today, NLP apps are everywhere. Siri and Alexa are examples of virtual assistants, Google Translate is an example of an automatic translator, chatbots are examples of customer service tools, semantic search engines are examples of search engines that look for meaning, automatic summarization systems are examples of systems that summarize information, and content generators are examples of systems that create content.

Modern NLP is going beyond just understanding text to include multimodal comprehension (language + pictures) and creative generating (storytelling, code creation, and music). There are still problems to solve, like data bias, understanding abstract ideas deeply, and dealing with false information that is automatically generated¹⁷.

1.3.4 Computer Vision

Computer vision is the branch of AI that lets systems get, interpret, and understand visual input like pictures or in the case of automated robots real time images. The end goal is to make it possible for computers to see and understand the world in a way that is alike to how people do.

In the past, computer vision relied on manual methods for finding things like edges, corners, and textures. But deep learning, especially convolutional neural networks (CNNs), has transformed the way things work: now, models can automatically learn important properties from raw data¹⁸.

Computer vision is used for a lot of things, such as:

- Recognizing faces (like unlocking your phone);
- Finding objects in moving surroundings (like self-driving cars);
- Medical imaging (like finding malignancies early);
- Automated quality control in manufacturing.

1.3.5 Autonomous Robotics

Autonomous robotics is a field that integrates skills from artificial intelligence, mechanics, advanced sensors and control theory, with the aim of designing robots capable of perceiving, planning and acting autonomously within complex and dynamic environments. An autonomous robot is able to perceive the surrounding environment through sensors such as cameras, lidars and sonars, build an internal map to locate itself and move effectively in space (through SLAM techniques - Simultaneous Localization and Mapping), plan optimal paths avoiding obstacles, and perform complex tasks without the need for direct human intervention.

The fundamental technologies that support autonomous robotics include motion planning, dynamic control, computer vision for interpreting the environment and reinforcement learning, which allows robots to continuously improve their operational strategies based on the experience acquired¹⁹.

1.4 Evolution of AI in the Industrial Context

1.4.1 Early Applications of AI in Industry (1980s–1990s)

Expert Systems were the most common use of artificial intelligence in industry environment in the 1980s and 1990s. These systems were made to help people make decisions about manufacturing by using logical rules that had already been set up. The most common uses were for setting up complicated items and organizing production tasks. In some circumstances, they helped figure out the best technological combinations to make a particular product when there were a lot of variables to think about and they all depended on each other. In other circumstances, they helped with the automatic planning of operational stages in factories by assigning resources, schedules, and orders based on certain principles. These systems were a real step in making industrial design and planning smarter by using "if-then" logic and cases that were already recognized.

1.4.2 Introduction of Machine Learning and Big Data (2000s–2010s)

Starting in the 2000s, the increasing capabilities in industrial data collection, thanks to the spread of IoT sensors, advanced ERP systems, and low-cost storage capacity, paved the way for the introduction of machine learning. By using Big Data and by having access to more computational power, AI models finally became capable of learning directly from experience,

overcoming the limitations of rigid rule-based systems. During that period the main application used in the industrial sector were predictive maintenance and automated quality control. The first one gave the ability to firms to anticipate failures by analysing weak signals in machine operation data. The second was able to detect products that do not comply with quality standards, thanks to the ability to find specific pattern in outgoing products. These innovations brought greater flexibility and responsiveness to industrial operations, reducing downtime, improving quality, and optimizing the use of production resources.

1.4.3 The birth of Smart Factories (2020s–Present)

The most recent chapter of AI application in the supply chain are Smart Factories. As already seen, many of the technologies that make Smart Factories "smart", such IoT, computer vision and collaborative robots ecc., were already deployed singularly in the industrial environment. The major change happened when artificial intelligence was added to integrate all these technologies together. Every technologies collect data that can be used not only from the original gatherer, but from every single autonomous system in the firm. These technologies also use deep learning and predictive analytics to turn the collected data it into real-time decisions that are autonomous, active, and based on the context. By combining data from sensors, ERP and supply chains, smart factories are now equipped with self-tuning and resilience capabilities. AI solutions automatically respond to variations in demand, logistics disruptions or production errors, making the production system more flexible and competitive²⁰. This evolution marks a turning point in the history of supply chain management, where AI no longer serves as a supportive tool but becomes a central orchestrator of operations. Smart Factories exemplify how AI can unify disparate technologies into an intelligent, adaptive, and interconnected ecosystem. As we look to the future, the integration of AI across the supply chain is set to deepen, enabling end-to-end visibility, real-time responsiveness, and sustainable optimization—paving the way for a truly autonomous supply network.

CHAPTER 2 – Supply Chain: Structure and Challenges

2.1 Introduction to Supply Chain

2.1.1 Definition

The supply chain represents the set of processes, resources and activities that are involved in the creation and distribution of a product or a service, from the initial supplier to the final customer. More precisely, it can be defined as an integrated network of organizations that collaborate to provide goods and services to markets. By doing so, this network has to manage simultaneously three flows: the physical flows of materials, information flows and financial flows.

Traditionally, the concept of supply chain has evolved from classic logistics, which was about the transportation and movement of goods. However, with the growing complexity of globalized markets, simple logistics management was no longer sufficient to guarantee competitiveness and efficiency. The need for a more systemic vision, capable of coordinating and integrating all the business functions involved in the creation of value, has therefore emerged.

To better understand what the Supply Chain is, we can analyse the discipline that manages it, namely Supply Chain Management. According to the definition proposed by the Council of Supply Chain Management Professionals (CSCMP), one of the main academic authorities on the subject, supply chain management includes "the planning and management of all activities involved in sourcing, procurement and conversion as well as all logistics management activities. It also includes coordination and collaboration with upstream and downstream partners in the chain, which may be suppliers, intermediaries, third-party service providers and customers."²¹

This definition highlights the importance of integration and inter-company collaboration as key elements for the success of the modern supply chain. In an increasingly dynamic market, the ability to act as part of a coordinated network becomes a critical element to address the challenges related to the volatility of demand, globalization, the pressure for greater sustainability but above all the political instability that emerged from the second decade of the 2000s.

2.1.2 Historical Evolution: From Traditional Logistics to the Modern Supply Chain

Until the 1980s, the work of logistics management was mainly focused on the operational efficiency of individual functions, such as transportation, storage and inventory management, often optimized in watertight compartments. Ballou described this functional approach in his academic manual "Business logistics/supply chain management"²²: he highlighted that the approach generated suboptimalities due to poor information sharing, high security costs and long transit times. Each department goal was to reach its own cost objective: however, by doing so they were reducing overall visibility and reactivity of the network.

With the maturation of the concept of Supply Chain Management, starting from the 1990s, the system moved to an "end-to-end" model that directly managed these inefficiencies. The most important innovation was the introduction of ERP systems and traceability platforms that made possible to integrate physical, information and financial flows in real time, reducing lead times and safety stocks thanks to shared visibility²³. Furthermore, the adoption of Sales & Operations Planning (S&OP) processes and simultaneous multi-level planning techniques has fostered a

strategic alignment between demand and production capacity, improving the agility in responding to demand peaks and resilience to disruption.

2.2 Supply Chain structure: stakeholders and main processes

2.2.1 Main Stakeholder

Within a well-structured supply chain, five main categories of actors can be distinguished, each with specific functions and responsibilities along the supply flow²⁴:

Suppliers

They are the first actors involved in the supply chain: they provide components or services to make the production process start. They can supply raw material (first level of suppliers), semi finished products (second level) or primary logistic services. To ensure quality, competitive cost and continuity of supply firms have to perform an effective supplier management, based on long-term relationships, performance assessments and collaborative practices.

Manufacturers

Then raw material are transformed into finished products through operational processes that include production planning, quality control, and plant maintenance, all done by manufacturers. Their main goals are to maximize plant efficiency and reduce setup times, integrate lean practices and digital technologies to respond quickly to changes in mix and volume.

Distributors

They act as intermediaries between manufacturers and retailers/end users. They manage warehouses that can be (centralized or regional), organize national and international transport and also offers value-added services like cross-docking, goods consolidation, co-loading. A good logistic network helps distributors to reduce transportation cost and lead time.

Retailer

They are the first point of contact with the consumer: they manage assortments, organize promotions and offers customer service. Inventory control at the point of sale and the ability to collect real-time sales data are essential to feed the replenishment and demand planning processes upstream.

End customers

They are the final destination of the products or services; their demand drives the entire chain. Feedback, purchasing behaviour and required service levels influence production, inventory and logistics decisions at all levels.

2.2.2 Main processes

Procurement and supply management

is the selection and purchase of raw materials, components and services from suppliers: the goal is to ensure continuity and quality at the best total cost. The main activities of this phase include supplier qualification, definition of sourcing strategies (i.e. deciding whether to source from one or more suppliers for the same purchasing category), contract negotiation and payment cycle: these choices are supported by the levels that suppliers can guarantee for various KPIs (key performance indicators) such as the procurement lead time and the Order Fill Rate (OFR, i.e. the percentage of satisfied orders).

Production and Operations Management

Manufacturing transforms purchased materials into finished goods by doing strategic planning and operational control. Key activities in this process include Material Requirements Planning (MRP) to calculate supply needs and times, production batch scheduling, preventive and predictive maintenance of plants, and monitoring Overall Equipment Effectiveness (OEE) to maximize resource utilization.

Logistics and distribution

This process coordinates inbound logistics and outbound logistics: the first activity deal with receiving and storing raw materials, instead the second deal with preparing and shipping finished products. When a company has to build its logistic network, it has to decide the optimal location of warehouses and distribution centres and it has to choose transportation methods (road, rail, sea, air) basing the decision on costs, times, and environmental constraints: at the end it has to apply Vehicle Routing Problem (VRP) models to optimize delivery routes. In some cases also reverse logistics has to be included in the building process of the network: it ensures the efficient collection of returns, reconditioning, and disposal, closing the product cycle.

Customer Service and After-Sales Management

Customer service ensures the connection between supply and demand by managing orders, deliveries and assistance requests. To monitor and evaluate customer satisfaction firms use the integrated Customer Relationship Management (CRM) systems: by using this technology companies can have access to costumers based indicators like Order Fill Rate (OFR), response time and Net Promoter Score (NPS). Instead, in the post-sales area, warranty, technical support, contractual maintenance and complaint management activities are essential for loyalty and corporate reputation. In the end, reverse logistics processes for returns and reconditioning, can transform a potential cost into an opportunity for upselling and continuous product improvement.

2.3 System for Supply Chain Management

2.3.1 Physical flow systems

The physical flow in the supply chain includes all the activities of movement, storage and shipping of materials, from receipt at the warehouse to delivery to the customer and collection of returns. To manage it efficiently, companies rely on a set of software systems and automation platforms, each dedicated to a specific phase of the journey of the goods.

Warehouse Management System (WMS)

The goal of a Warehouse Management System (WMS) is to manage heterogeneous warehouse activities by coordinating them through information on incoming goods flows, outgoing flows, and quantities in stock, thus providing the ability to monitor logistics processes in real time. A good WMS connects to technologies such as barcode scanning, RFID, mobile devices, robotics, and augmented reality, and integrates with external systems such as ERP, TMS, and logistics software. These integrations improve accuracy, operational efficiency, and reduce errors, meeting the need for speed and control in the modern supply chain context²⁵.

Transportation Management System (TMS)

It deals with the planning and monitoring of both inbound and outbound transportation. The TMS solves the Vehicle Routing Problem, selects the most cost-effective carriers, manages delivery windows and provides dashboards for real-time tracking of shipments, ensuring that products leave the warehouse according to the established times and costs²⁶.

Package Sizing and Optimization Systems

They are used in the packaging phase: by analysing volumetric weight, ideal pallet configuration and optimal number of packages, the system maximizes space, reduce transportation costs and helps achieve the goal of "transporting as little air as possible". Thanks to 3D scanners and pallet optimization software, they automatically update shipping parameters and feed the WMS with up-to-date data on volumes and dimensions.

2.3.2 Information flow systems

The information flow aims to ensure proper sharing of information that typically arises from separate processes within the supply chain. This ensures that all operational decisions, from demand forecasting to shipment tracking, are based on up-to-the-minute data. To orchestrate this data along the supply chain, companies adopt various specialized systems.

Enterprise Resource Planning (ERP)

ERP is the platform where information on past sales data, inventory, production plans and future market demand converge. SAP is one of the most widely used systems today, ensuring that multinationals have the right consistency of information across departments and automatically synchronizing processes. An example of a user of the system is the Italian multinational Barilla G. e R. Fratelli S.p.A and in particular the Production department. Through the ERP, planners can know in real time the stock status in the various warehouses of the network and in particular in those of the production plants. By associating this information with the future demand data developed by the Demand Planning office, employees can develop a Production Schedule for each Plant in the network, specifying which item to produce and the period within which it should be available. Then this information is shared at the Plant itself, where a "Process Follower" sees the scheduling in the ERP and prepares a raw material and operator allocation plan for each line affected by the production schedule.

Electronic Data Interchange (EDI) e Application Programming Interface (API)

Application Electronic Data Interchange (EDI) and Application Programming Interface (API) are two approaches used to exchange information between companies along the supply chain. EDI is a historic and widely used technology that allows business documents such as orders, invoices or shipping notices to be transmitted in a standard, structured format. It typically works in "batch" mode, that is, collecting and sending groups of documents at regular intervals, ensuring consistency, traceability and compliance with business partner requirements.

APIs, on the other hand, are modern interfaces that enable communication between systems in real time. With APIs, information can flow instantaneously between suppliers, customers and management systems (ERP, TMS or WMS), facilitating visibility and responsiveness along the supply chain. However, they require more coordination in developing and managing connections. In many companies, EDI and APIs do not compete, but are used together: EDI for massive, standardized exchange of official documents, APIs for fast updates and flexible integrations that improve collaboration and operational efficiency²⁷.

Business Intelligence e Dashboarding

The definition of business intelligence includes all those processes and tools that are used to collect, analyse and transform business data, and then use them to make operational or strategic data-driven decisions. Business intelligence allows for Dashboarding, i.e., creating interactive dashboards that allow for immediate visualization and, more importantly, understanding of raw business data, as well as Key Performance Indicators, i.e., performance measures against business objectives. BI dashboards enable the aggregation of data from various databases, their real-time updating, and the creation of customizable views based on the interests of

stakeholders. Dashboards, as opposed to traditional reports, are made to promote interactive and exploratory analysis through drill-downs, dynamic charts, and filters. This method promotes ongoing performance monitoring, enhances comprehension of operational and strategic trends, and helps the organization make better decisions²⁸.

2.3.3 Financial flow systems

Financial flow coordinates payments, invoicing, credit management, and working capital optimization among all supply chain stakeholders. To support this, companies adopt dedicated platforms:

Electronic Invoice Presentment & Payment (EIPP)

One of the most used solution to manage the invoice cycle is the Electronic Invoice Presentment and Payment (EIPP): this a digital solution automates the entire invoicing cycle-from the creation and electronic submission of invoices to online payment by the customer. These systems allow firm's software to send autonomously invoices via web or email portals: by doing so, customers are enabled to view and settle documents securely and quickly. Through integration with ERP software, EIPP automatically generates invoices, applies validity checks, and facilitates matching of payment and document. Benefits include paper reduction, accelerated payment times, fewer data entry errors, improved cash flow management, and increased customer satisfaction²⁹.

Treasury Management System (TMS)

Treasury management involves managing a company's cash flows and financial decisions, providing governance over liquidity, credit line maintenance, investment optimization, and fund use. A Treasury Management System (TMS) automates the treasury management process, providing better visibility into cash and liquidity, increased control over bank accounts, compliance standards, and better financial transaction management. TMS also offer real-time financial information, simplifying reporting and cash forecasting. They can automate payments, especially for large transaction volumes, and streamline reconciliation processes. Modern Treasury's Payments product offers custom payment controls for secure payments. A good TMS can streamline reconciliation by matching payments with transactions in bank statements, reducing manual work and errors³⁰.

2.4 Current challenges

2.4.1 Demand volatility

A major source of complexity in supply chain management is demand volatility: it's mostly caused by customer purchasing decisions that are often unpredictable variables for business systems. This instability is amplified along the supply chain through the phenomenon known as the bullwhip effect, where little changes in demand downstream (at the end consumer) cause disproportionate fluctuations in orders placed by upstream actors (wholesalers, manufacturers, suppliers)³¹. Under conditions of high volatility, traditional planning systems based on historical data become less effective, forcing firms to have high buffers of inventory, extra production capacity or conservative sourcing strategies: all of these methods prevent out of stock but have negative consequences on costs and operational efficiency. Volatility can be caused by cyclical factors (such as seasonality and consumer trends) or by sudden exogenous events (natural disasters, regulatory changes, geopolitical crises), which make demand behaviour not only unstable but also difficult to predict with conventional forecasting methods.

In recent years, demand volatility has taken on even more difficult characteristics for companies to manage. The paradigm began to shift with the pandemic from COVID-19: travel

disruption caused demand spikes in some sectors (e-commerce being one example) and paralysis in others (such as automotive) highlighting the fragility of global supply chains.

Before the pandemic, globalization has made national supply chains even more interconnected, increasing the phenomenon of labour outsourcing, especially in Southeast Asian regions. A single event in one region can now generate ripple effects on the entire global supply network, amplifying fluctuations in supply and demand.

One of the latest game-changers, have been social networks: platforms such as Instagram and Tik-Tok can generate unpredictable viral trends that can quickly influence consumer preferences, especially among young people. This phenomenon has been observed in various sectors, from fashion to food, where so-called “microtrends” can emerge and spread very quickly, testing the responsiveness of supply chains. It should be noted that these platforms are increasingly becoming showcases, where companies can show their products, often running the risk of overselling. Tik-Tok itself states that 70% of its users find new brands on its platform. Additionally, 75% of users are likely to make a purchase while using TikTok. Ultimately, 83% of users report that TikTok influences their purchasing decisions³².

To address all these new challenges, companies are adopting demand sensing tools and predictive techniques based on Big Data and Machine Learning, which integrate real-time sales data, social media signals, weather data, and economic information to continuously update forecasts and make the supply chain more agile and responsive.

2.4.2 Environmental and social sustainability

The concept of a sustainable supply chain stems from the desire of companies to combine their economic goals with environmental and social goals. In order to design, manage and optimize a supply chain with sustainability, it is necessary to persevere in each of these steps the goal of reducing environmental impact through practices such as energy efficiency, recycling and waste minimization: in addition, respect for human rights must not be overlooked on any continent the supply chain operates, ensuring decent working conditions for its employees and monitoring the respect its partners have for them³³. This transition to a sustainable supply chain has been necessitated by a combination of factors. The first is the growing expectation that investors, but especially consumers, have in taking advantage of products that reflect a lifestyle increasingly focused on respect for the environment and disadvantaged populations. Second, institutions have also begun to move under pressure of consumers, increasing regulatory pressure on companies. One example is the European Union's decision to ban the production of carbon dioxide-emitting cars from 2035³⁴. This decision has severely challenged the automotive industry, prompting companies to rethink their long-term strategies: the result has been a radical overhaul of corporate product portfolios and the adoption of more environmentally focused communication.

The main challenges that emerge in building a sustainable supply chain are:

- **Material traceability:** it is essential to be able to trace the path of materials throughout the chain-from raw material to finished product-to ensure environmental and social compliance. If companies don't monitor the origin of raw material used, the probability of involvement in unethical practices like deforestation , child exploitation or use of uncertified materials, increases.
- **Responsible supplier management:** because a supply chain may involve hundreds or thousands of suppliers at different levels, it is critical to establish clear selection and evaluation criteria based on ethical and environmental standards. This requires regular

audits, supply contracts with ESG (Environmental, Social and Governance) clauses, and training programs for at-risk partners.

- Reducing carbon footprint: a significant share of global gas emission is caused by manufacturing, transportation and logistics processes: by optimizing delivery routes, adopting electric vehicles in the companies fleet and improving the energy efficiency of facilities, are key actions to reduce overall environmental impact.
- Circular economy: is essential to remove the old paradigms of “take-make-dispose”. Firms have to redesigning products and processes to encourage reuse, recycling and recovery of materials at the end of life. By doing so it will be generated a new circular economy that will reduce dependence on virgin resources and limits waste accumulation, contributing also to global sustainability goals.

At the end, it should be emphasized that the recent improvements that companies have made in the matter of sustainability and respect for rights, are due primarily to the efforts of consumers. Without a public aware to these issues and able to exert pressure demanding constant improvement, none of this would be possible; instead, companies would have no real incentive to review harmful practices and invest in more ethical and responsible solutions.

2.4.3 Risk management and resilience

Risk management in the supply chain is finding, evaluating, and reducing the chances of events that could disrupt the movement of goods, information, and money along the supply chain. A supply chain that is strong can not only handle shocks, but it can also change swiftly and get back to normal quickly.

Because of worsening global geopolitical balances, managing risk in the supply chain has become even more important in 2024 and 2025. Recent events have shown how easily political, economic, and logistical shocks can affect worldwide supply chains.

War in Ukraine and the European energy crisis

The continuation of the war in Ukraine has had a systemic impact on global supply networks, especially those for energy and food. Because Russia's natural gas and oil supplies have been cut off, Europe has had to quickly change the way it gets its energy. It is now getting more LNG from the US and Qatar and investing more quickly in renewable energy. From the point of view of the industrial supply chain, a lack of raw materials like steel, fertilizer, and wheat has slowed down output and raised prices in important industries like chemicals, automotive, and agro. Also, the closing of numerous land and rail routes has made it harder for sea and air freight to get things done, which has made international deliveries take longer on average. The scenario has made European corporations look for more suppliers, which means they are less dependent on unstable areas and are putting money into making their most important products more regional³⁵.

U.S.-China Trade Conflict

In the 2024-2025 biennium, the economic competition between the US and China has quickly gotten worse. The US has put export limits on the sale of advanced technologies to China: this ban include semiconductors, microchips and artificial intelligence. This decision affects not only Chinese suppliers but also worldwide production chains that depend on these parts and perform part of their operation in the Est. As a result, also China has put limits on the export of important resources like gallium, germanium, and rare earths. These elements are necessary for the technology, telecommunications, and green energy industries: also they are mainly mined in the Asian country, thus becoming a scarce resource³⁶. This has put a lot of stress on the high-

tech industries, which has slowed down the making of electric vehicles, electronic devices, and photovoltaic systems. Because trade cooperation is getting worse, many corporations are changing how they source goods. They are prioritizing friendshoring (putting production in countries that are politically aligned) and speeding up technology reshoring initiatives, especially in the US, Europe, and Japan³⁷.

Logistics crisis in the Red Sea and Taiwan Strait.

Shipping in the Red Sea and Taiwan Strait has experienced a phase of high instability.

Attacks on merchant ships by armed groups in the Red Sea region (linked to tensions in Yemen and the Horn of Africa) have forced many shipping companies to divert routes by circumnavigating Africa via the Cape of Good Hope, resulting in increased delivery times between Asia and Europe³⁸.

At the same time, the Taiwan Strait, has been the scene of military manoeuvres and diplomatic tensions between China, Taiwan and the United States. The threat of a naval blockade or sanctions has led many logistics operators to seek safer, albeit more expensive, alternatives, negatively affecting the fluidity of routes between Asia and North America³⁹. The impact on global supply chains was immediate: increased container freight rates, delivery delays, shortages of critical components (especially electronics), and the need to renegotiate logistics and insurance contracts on more onerous terms.

2.5 Summary and outlook

This chapter showed that the modern supply chain is a dynamic system: it is made up of a network of connected players who coordinate the flow of goods, information, and money. The shift from traditional logistics to an end-to-end approach has turned the supply chain into a real competitive advantage for companies. Over the last few decades, innovations like ERP systems for process integration, WMS and TMS for logistics automation, and Business Intelligence platforms for data visualization and analysis have improved visibility and lowered costs, making supply chain management more flexible and customer-oriented.

Even with all these improvements, the global context is still complex and full of uncertainty. Things like demand volatility, made worse by the bullwhip effect, geopolitical tensions, and increasing regulations on environmental and social sustainability mean that supply chain models need to keep evolving. Building resilience, meaning the ability to absorb shocks and get things back on track quickly, has become a key priority. This calls for strategies like diversifying suppliers, using friendshoring or reshoring, improving raw material traceability, and adopting circular economy practices.

Looking ahead, technology will play a big role in dealing with these tough and hard challenges. Tools based on Artificial Intelligence and Machine Learning, like real-time demand sensing, network design optimization, predictive models for maintenance, and anomaly detection for risk management, will be essential to make forecasting and decision-making much better. Using digital twins and advanced simulations will also let companies try out different what-if scenarios, which helps them plan smarter and react faster when things don't go as expected.

Looking ahead, the main challenge will be finding the right balance between efficiency, flexibility, and sustainability. A supply chain will need to be not only high-performing, but also ethical, transparent, and resilient. Companies that manage to use new technologies to build smarter, more collaborative, and more sustainable supply chains will be the ones able to gain a lasting competitive advantage in a constantly changing market.

CHAPTER 3 – AI's applications in Supply Chain

3.1 Introduction

3.1.1 Chapter goals

In recent years, supply chains have faced an increasingly unpredictable environment, characterized by unforeseen events (pandemics, geopolitical tensions, energy crises) and growing management complexity due to globalization, the personalization of demand, and pressure on sustainability. Faced with these challenges, it has become clear that traditional management models based on historical data, rigid planning, and slow reactions are no longer sufficient.

In this scenario, artificial intelligence (AI) emerges as a strategic lever for the digital transformation of value chains. Thanks to techniques such as machine learning, natural language processing, and predictive analytics, AI is able to process large volumes of heterogeneous data, detect subtle patterns, and support automated or semi-autonomous decisions across all levels of the supply chain. The practical applications of these technologies are now widespread and documented in numerous industrial cases, with measurable effects in terms of cost reduction, increased agility, and improved service levels. To provide a numerical representation of these benefits, in early 2024, the Georgetown Journal of International Affairs, selecting a sample of "AI early adopters," demonstrated how this technology can reduce logistics costs by 15%, improve inventory levels by 35%, and enhance service levels by 65%⁴⁰.

Having understood the potential challenges and benefits of AI in the supply chain, this chapter aims to illustrate its main applications, dividing them into 4 thematic areas:

- Demand forecasting,
- Inventory optimization,
- Predictive logistics,
- Quality control,

The chapter can be seen as a pivot to move from the management challenges seen in chapter 2 and the AI solution discussed in chapter 4 through real world case studies.

3.1.2 Types of Artificial Intelligence in the Supply Chain

Types of AI that can be used in Supply Chain are classified in three main categories: analytical, predictive, and prescriptive. The classification helps to understand what each sort of AI can perform to help the decision making process⁴¹.

The first type is Analytical AI, which is used to analyse historical data to find trends, outliers, or origin of variation in performance indicators. This AI doesn't give an answers but instead, it helps to comprehend what happened and how it affected firm's operations.

Predictive AI employs machine learning and statistical models to understand what might happen in the short or even medium-longer future of firms operation. Based on historical data, trends, and outside factors, it is commonly used for demand forecasting, predictive maintenance. It makes forecasts more accurate, lowers safety inventories, and makes the best use of manufacturing capacity in the supply chain.

Prescriptive AI is the most advanced type since it directly tells companies what to do to reach a specific goal, like lowering expenses or raising service levels. This kind of AI combines predictive models with optimization algorithms to give clear advice or even start automatic

actions. It can be used in the supply chain to make flexible plans for manufacturing and distribution, assign transportation loads in real time, or move resources around, based on what the business needs in a specific moment.

These three types are not mutually exclusive, but rather complementary: a mature digital supply chain uses descriptive analytics to monitor, predictive analytics to anticipate, and prescriptive.

3.2 Demand Forecasting

3.2.1 Definition

Demand forecasting is one of the most important processes within a manufacturing company, as it attempts to estimate in advance the quantity of products or services required by customers within a specific time frame. Accurate estimates will enable other business units to successfully manage their operations. Effective demand forecasting allows for effective planning of production and logistics flows, while reducing waste, delays, and unnecessary costs. Furthermore, as mentioned in the previous chapter, having flexible demand forecasting processes allows for adapting to customer purchasing behaviours, which are increasingly driven by unpredictability and decisions dictated by factors beyond the company's control.

3.2.2 Predictive Models vs. Traditional Approaches

Traditionally, companies have relied on classic statistical methods to estimate future demand behaviour: the most common ones were linear regression, ARIMA model and Croston model. These models are relatively simple to implement, easy to interpret, and require a limited amount of historical data.

Linear Regression

Among the most common methods for demand forecasting, linear regression is one of the simplest and most historically used models. This approach aims to explain the behaviour of the dependent variable (for instance future demand) as a linear function of one or more independent variables, such as past sales, seasonal trends, prices, or macroeconomic indicators. In mathematical terms, linear regression is expressed by the formula in the following image.

The diagram shows the linear regression formula $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \varepsilon$. Arrows point from descriptive labels to the corresponding parts of the formula: 'Dependent Variable (Response Variable)' points to Y ; 'Independent Variables (Predictors)' points to the X terms; 'Y intercept' points to β_0 ; 'Slope Coefficient' points to β_1 and β_2 ; and 'Error Term' points to ε .

Where:

- Y is the variable to estimate (for instance future demand),
- $X_1, X_2, \dots, X_n$ are the explanatory variables (e.g. past sales, discounts, seasonality, etc.),
- β_0 is the intercept: it's the value of Y when all the X s are equal zero.
- $\beta_1, \dots, \beta_n$ are the coefficients that measure the effects of each X on the Y variable,
- ε is the error term, which represents the component not explained by the model.

The underlying principle of linear regression is quite simple: it starts from the idea that there is a stable and proportional relationship between each explanatory variable and demand. In

other words, if a variable X increases Y is also expected to increase or decrease accordingly to the estimated coefficient.

But this paradigm does have some relevant problems. It assumes that each variable influences demand independently and with a constant impact over time. It is also very sensitive to multicollinearity, which is when two or more independent variables are correlated to each other. The model may not work right if the variables are multicollinear. This is a big difficulty when trying to put theory into practice. In the supply chain, demand does not always follow linear and orderly patterns: it can depend on the combined effects of multiple variables, unexpected events, or trends that change over time. These are dynamics that a linear model, by its nature, struggles to capture.

For this reason, while useful in simple or very stable contexts, linear regression shows its limitations in complex scenarios such as those of modern supply chains, where uncertainty, volatility, and strong interconnections between factors reign.

ARIMA Model

One of the most widely used methods for demand forecasting is the Autoregressive Integrated Moving Average model. Using this method, a historical series such as market demand can be analysed and, after a series of processes, the new component of the series, the forecasted value, can be obtained.

This model combines three main blocks:

AR - Autoregressive: In this module, it is assumed that future demand is linked to demand already recorded over time. Linear relationships between the variables are therefore sought, just as in linear regression.

I - Integrated: This component of the model aims to stationarise the time series, eliminating trends and seasonality that could compromise accurate forecasting. This is done by calculating the differences between consecutive observations (for example demand of the current month minus that of the previous month), thus analysing variations rather than absolute levels.

MA - Moving Average: This takes into account the forecast error made in previous periods, assuming that the error is not random but has a systematic structure. In practice, the model corrects the current forecast based on deviations observed in the past.

The ARIMA model is parameterized with three integer values (p, d, q), where:

- p indicates the order of the AR component (how many past values to consider),
- d indicates the number of differencing needed to make the series stationary,
- q indicates the order of the MA component (how many past errors to include).

The three components are then combined in the following equation:

$$y'_t = c + \underbrace{\varphi_1 y'_{t-1} + \dots + \varphi_p y'_{t-p}}_{\text{lagged values}} + \underbrace{\theta_1 \varepsilon_{t-1} + \dots + \theta_q \varepsilon_{t-q} + \varepsilon_t}_{\text{lagged errors}}$$

intercept
lagged values
lagged errors

differenced time series

In short, ARIMA tries to guess the next value in a time series by looking at prior values, historical forecast mistakes, and any trends that show up in the data. It gives each one a weight based on how important it is. The model uses the historical series, makes some changes, such as differencing to get rid of patterns, and then predicts the next point based on what has happened previously and how wrong the last forecasts were.

It's a sensible way to do things, and it usually works well when the data is stable and follows a pattern that can be seen.

That being said, ARIMA doesn't always work well when things are more complicated. It can be hard when the data is very volatile or affected by outside forces, like demand in modern supply chains. It doesn't fully fail, but it does become less reliable when things get too crazy or out of control⁴².

Croston Model

In supply chain management, one of the most persistent challenges is forecasting intermittent demand, demand marked by long stretches with no orders, interrupted by sudden and irregular requests. Situations of this kind often arise in areas such as spare parts, seasonal goods, or slow-moving products. Standard techniques like linear regression or ARIMA are usually unsuited to the task, since they treat the sequence of zeros and positive values as though it followed a linear pattern, which in practice leads to systematic over- or underestimation. To confront this problem, J.D. Croston introduced in 1972 a method tailored to such cases: the Croston model, which has since become a cornerstone in many inventory management systems. Its key contribution is the idea of disentangling the estimation of demand size from that of demand frequency, thereby overcoming the weaknesses of traditional approaches⁴³.

The Croston model can be expressed in a compact way as in figure:

$$\text{if } d_t > 0 \begin{cases} a_{t+1} = \alpha d_t + (1 - \alpha)a_t \\ p_{t+1} = \alpha q + (1 - \alpha)p_t \\ f_{t+1} = \frac{a_t}{p_t} \end{cases}$$

$$\text{if } d_t = 0 \begin{cases} a_{t+1} = a_t \\ p_{t+1} = p_t \\ f_{t+1} = f_t \end{cases}$$

Where:

- d_t is the observed demand at time t ,
- a_t is the updated estimate of the average quantity of demand (i.e., how much is ordered on average each time a demand occurs),
- p_t is the periodicity, i.e., the average interval between two positive demands,
- f_t is the demand forecast for each period,
- $\alpha \in (0,1)$ is the smoothing factor,
- q is the number of periods since the last non-zero demand.

This formulation makes explicit the conditional behaviour of the model: when a positive demand is observed, the averages are updated with the new value and the forecast is recalculated; when demand is zero, no parameter is updated and the forecast remains unchanged⁴⁴. Still, the model has some important shortcomings. It tends to overestimate average demand, since the parameters remain fixed during periods without orders. It also

struggles to account for elements such as trend or seasonality, and it does not easily adjust when demand undergoes structural changes or is influenced by external shocks.

In the end, although it is simple and often effective in dealing with intermittent demand and in describing relatively stable scenarios, Croston's model shows structural limits that reduce its usefulness in more dynamic environments.

3.2.3 Machine Learning models

In recent years, the spread of artificial intelligence, especially machine learning and deep learning, has introduced new approaches capable of addressing many of these limitations. The following section will look at predictive models such as gradient boosting and artificial neural networks, which learn directly from the data without requiring the analyst to define in advance how the variables relate to one another. They can also adapt as new information becomes available and are able to work with very large datasets, including those that are only partially structured. Beyond that, AI makes it possible to uncover patterns that would be hard to detect with classical techniques, to react more quickly to shifts in the surrounding context, and to draw on a wide range of data sources, from weather records to social media activity, all the way to exchange rate movements.. At the same time, however, their use also depends on certain prerequisites, such as data quality, sufficient computational resources, and solid analytical expertise. In short, AI-based predictive models do not entirely replace traditional methods, but they extend and strengthen them, providing higher accuracy and flexibility in contexts marked by uncertainty.

Gradient Boosting Machines (GBM)

Gradient Boosting Machines (GBMs) are often regarded as a very effective approach for problems such as regression and classification, and in recent years they have become especially popular in industry for demand forecasting. The idea is fairly straightforward: instead of relying on a single model, a GBM combines many weak learners, usually decision trees, so that, together, they form a much stronger predictor.

A decision tree itself can be pictured as a flow of questions and answers: each inner node tests a variable, the branches represent the possible conditions, and the leaves give the final output. On its own, however, a single tree that is not sufficiently deep has limited predictive power. This is where the idea of building collections of trees comes in.

Random Forests offer one way of combining trees. Random Forests make use of a technique known as bagging. In practice, this means that many trees are built at the same time, each on a different random subset of the data, and the final prediction comes from combining their outputs, for example by averaging the results or taking the majority vote. This technique lowers variance and usually makes the model more stable, though it tends to produce predictions that are somewhat averaged out, without directly targeting the errors.

Gradient Boosting takes a different path. Instead of growing trees at the same time, it builds them one after another, and each new tree is trained to improve on the shortcomings of the previous ones. In this way, the model directs more attention to the hardest observations to predict. This sequential refinement gives the final model a very high predictive capacity, though at the cost of greater computational complexity⁴⁵.

In the context of weekly demand forecasting for a retail product, GBM can be explained through a sequence of steps. The process begins with a very rough estimate of the target variable, for instance, the number of units sold in a given week. This first guess is usually taken

as the average observed in the training data and acts as a starting point that the model will refine step by step.

From here, the model calculates the residuals, that is, the difference between the actual demand recorded in the data and the forecast produced up to that point. In practice, this step amounts to checking how far the predicted weekly sales are from the actual figures recorded in the data. A decision tree is then trained to look for recurring patterns in those differences, in order to explain the source of the error. For instance, the tree may point out that demand systematically increases during promotional periods or that unusual weather conditions tend to alter sales levels. The adjustments suggested in this way are introduced gradually, by means of the learning rate, a parameter that regulates the contribution of each new tree. This ensures that forecasts are corrected step by step, without producing abrupt or implausible jumps.

This cycle, calculating residuals, fitting a new tree, and updating the prediction, is repeated several times. With each round, the model puts more weight on the observations that remain hardest to predict. Once the process has run for a fixed number of iterations, or when an early stopping condition is triggered, the outcome is a final model that brings together all the trees, each weighted according to its contribution.

The advantages that GBM offers over traditional models are several. First of all, it performs very well in modelling non-linear relationships between variables, a crucial aspect in demand forecasting, since sales are often shaped by complex and interdependent factors such as seasonality, promotions, macroeconomic indicators, or external data sources like weather conditions, special events, and online trends. In contrast with more traditional linear models, GBM manages to capture these dynamics in a more flexible way, even when different features interact with each other.

An additional important feature of GBM is that it can rank the importance of the input variables, showing which ones have the greatest effect on the forecast. This makes it easier for supply chain managers to understand the main factors driving demand and to plan targeted actions, for instance by increasing promotional efforts in periods of higher sensitivity. In applied settings, recent versions of GBM such as XGBoost, LightGBM, and CatBoost are often preferred because they scale well. These implementations can handle very large datasets while still producing accurate results within acceptable computation times. With such tools training can be completed on millions of observations in a short period, without sacrificing efficiency or predictive accuracy.

Another advantage is that GBM can also manage missing values and categorical variables, which are common in real-world datasets, without requiring complicated preprocessing steps. This reduces the effort needed in the preparation stage and allows the model to be integrated into company information systems more smoothly.

To sum up, Gradient Boosting Machines can be considered among the most effective tools for demand forecasting in supply chain management. Their strength lies in the ability to learn complex links between variables, to cope with heterogeneous datasets, and to provide forecasts with a high level of accuracy. What makes them appealing in a business setting is the balance they offer between accuracy, flexibility in use, and the possibility of interpreting the results. As a result, GBMs turn out to be especially useful in business practice, where accurate forecasts can support better decisions and contribute to smoother and more efficient logistics.

Support Vector Regression (SVR)

Support Vector Regression (SVR) is a supervised learning technique derived from the Support Vector Machine (SVM) model, which was introduced by Vapnik and Cortes in the 1990s as a method for both linear and non-linear classification. Suppose we are given a labelled dataset divided into two classes. The goal of an SVM is to identify a hyperplane, expressed as a linear function $f(x) = w^t x + b$, that separates the two classes in such a way that the margin, the distance between the hyperplane and the closest points from either class known as support vectors, is maximized. In practical terms, this means creating a rule that can assign new observations to one class or the other. Given a new point x , the model checks on which side of the hyperplane it lies and classifies it accordingly. Choosing the hyperplane that maximizes the margin makes the model more robust to noise and small changes in the data. In this way, the goal is not only to match the training set but also to achieve a solution that generalizes well⁴⁶.

This approach was later extended to regression problems, leading to what is known as Support Vector Regression (SVR). In this case, the aim is not to separate classes, but to identify a function that stays as close as possible to the observed data, within a tolerance level ε . A distinctive feature of SVR is the introduction of an ε -insensitive margin, where deviations smaller than ε are not considered as errors. The optimization task is therefore to find a function $f(x)$ that combines low structural complexity, meaning a wide margin, with respect for the error constraints. In demand forecasting, SVR is applied to capture the link between future demand (the target) and explanatory variables such as price, promotions, seasonal factors, weather conditions, and past sales patterns. Each observation in the dataset can be written as a vector $x = (x_1, x_2, x_3, \dots, x_n)$, where every element x_i represents an independent variable. The purpose of the model is to learn a function $f(x)$ that produces, for any given input vector, a prediction of demand y . What matters is that the prediction does not depend on a single factor, but on the joint effect of many variables considered together. When the model is trained, SVR does not impose any predetermined shape on the link between inputs and outputs. It infers this link directly from the data, learning how each variable x_i influences the forecast. For example, if the dataset shows that promotions are usually followed by higher sales, the algorithm will capture this by giving more importance to that variable. Conversely, a higher price may emerge as a factor associated with lower demand^{47,48}.

Thanks to its ability to generalize and to remain robust in the presence of noise, SVR can be considered a solid option in situations where the amount of data is not very large but variability is significant.

3.2.4 Deep Learning models

Artificial Neural Network (ANN)

Many researchers describe these networks as a bridge between machine learning and deep learning. This is true if we look at theory, but also if we think about how they are applied in practice. The basic idea comes from the brain, although only in a very rough way. A network is formed by a lot of small units, artificial neurons, which are arranged into layers. Normally there is an input layer first, then one or more hidden layers in the middle, and at the end the output layer. A neuron just gets some information, does a quick operation on it, and hands it on to the next ones. Step after step, this simple routine allows the network to build up its final prediction.

In methodological terms, ANNs are usually included among supervised learning methods, since they rely on labelled data to learn how to associate inputs with the correct outputs. What has changed the most in recent years is the growing depth of these networks. By adding more

hidden layers and by relying on much stronger computational resources, researchers have been able to extend the potential of ANNs, and this is basically how deep learning was born. In practice, deep learning makes use of very deep neural networks that can capture complex and abstract patterns in the data, often reaching results that are beyond what traditional statistical models or earlier learning methods could do⁴⁹.

The origins of artificial neural networks can be traced back to the 1940s, with the formal model introduced by McCulloch and Pitts (1943). In their work, a neuron was described as a logical unit that could activate depending on input thresholds. A first practical development came later with Frank Rosenblatt's perceptron (1958), a single-layer neural network able to learn weights for classifying linearly separable data. However, the enthusiasm around neural networks slowed down considerably after the publication of Minsky and Papert's book (1969), which highlighted their theoretical limits. In particular, they showed that the perceptron was unable to solve problems that were not linearly separable, such as the well-known case of the XOR logical operator⁵⁰.

Interest in artificial neural networks (ANNs) started to grow again during the 1980s, mainly because of the work of David Rumelhart, Geoffrey Hinton, and Ronald Williams. In 1986 they published a very influential paper where they introduced the backpropagation algorithm. With this method it finally became possible to train, in an effective way, neural networks made of several layers, which until then had remained more of a theoretical idea than a practical tool. The key challenge was that, even if researchers already knew how to update the parameters (the so-called weights) of very simple models like the perceptron, there was still no reliable way to modify the weights of the hidden layers inside more complex networks. These hidden layers do not produce a direct output, and for this reason it was unclear how to modify their behaviour according to the overall error of the network.

The backpropagation algorithm solved this issue by creating a method to send the error signal "backwards" through the network. It begins from the difference between what the network predicts and the real value (that is, the error), and then it estimates how much every single neuron, in each layer, has contributed to that error. Based on this information, the weights are updated. The same procedure is repeated over and over, until the total error of the network goes down below a certain level. Thanks to this innovation, it became possible to actually build and train multilayer networks able to face more complex and non-linear problems. In other words, problems where the link between the input variables and the output is not direct or simple at all. This result, which was at the same time technical and also conceptual, opened the road for the creation of deeper networks and, more generally, for what we call today modern deep learning⁵¹.

The word deep is used because of the presence of several hidden layers, which give the network the possibility to learn hierarchical and more and more abstract representations of the input data. In other words, an ANN with only one or two hidden layers can still be considered a "classical" machine learning model, while a deeper network, with many layers, more complex activation functions, together with normalization and regularization techniques, clearly belongs to the field of deep learning.

From a structural point of view, an artificial neural network is made of elementary computational units called artificial neurons, which are loosely inspired by the functioning of biological neurons. The neurons are put together in layers: first the input layer, that receives the variables from the dataset; then one or more hidden layers, where the information is really worked out; and finally the output layer, which gives the prediction.

Each neuron receives as input a set of numerical values coming from the neurons in the previous layer, combines them linearly by means of synaptic weights (w) and a bias (b), and then applies a non-linear activation function (for instance ReLU, sigmoid or tanh) to obtain the neuron's output:

$$a = \phi(\sum w_i x_i + b)$$

where ϕ is the activation function, x_i are the inputs, and w_i the corresponding weights. As explained before, the learning process is carried out by an iterative algorithm called backpropagation, usually in combination with numerical optimization methods such as gradient descent. During training, the model calculates the prediction error (for example, the squared difference between the observed and the predicted value) and updates the weights throughout the network by propagating the gradient of the error backwards, in the opposite direction of the data flow. This process goes on until the error reaches a minimum threshold, or until the maximum number of training epochs is completed.

The presence of several hidden layers makes it possible for the network to learn progressively more abstract and hierarchical representations of the input data, giving ANNs the ability to model highly non-linear relations and to capture complex interactions among variables, things that usually escape traditional models. This particular feature of ANNs makes them a powerful and flexible solution for demand forecasting in the supply chain, especially as an answer to the structural limits of traditional models. Unlike those models, which often assume linear relations or rely on strong hypotheses about the data (such as stationarity in time series), ANNs are able to learn the underlying relationships directly from the data, without forcing a rigid structure in advance. ANNs are trained on the same historical datasets used by traditional approaches, with independent variables that usually include past sales, product prices, the presence or absence of promotions, seasonality indicators (such as month, day of the week, or holidays), external factors (weather conditions, local events, special openings), as well as stock levels and other logistical information. What comes out of the network is the predicted demand for the next period. Thanks to the fact that they have many layers and non-linear activation functions, ANNs can also capture complicated interactions among variables, like a promotion being useful only in certain months, or the effect of price changing with the sales channel.

One of the main advantages of ANNs is their ability to generalize: once they are trained on a representative dataset, they can forecast demand even in scenarios that are not exactly the same as the ones observed, but similar in structure. This aspect is particularly useful in industries characterized by strong volatility, such as fast-moving consumer goods, fashion, or e-commerce. Finally, ANNs can be integrated into automated forecasting pipelines, updating themselves regularly with new data. In many companies, they are already used as part of advanced demand sensing systems, which can quickly react to changes in consumer behavior, improving forecasting accuracy and optimizing inventory management as well as operational planning⁵².

3.3 Inventory optimization with AI

3.3.1 Definition and Challenges

Optimizing inventory is often one of the hardest tasks when trying to keep a supply chain efficient. Well-planned stock levels can help cut operating costs while still allowing the company to meet customer demand. In practice, the goal is to set, for every product and at

every stage of the network, inventory levels that balance product availability with the overall cost of logistics.

Inventory doesn't serve a single purpose, and it's not all the same. In practice, it can be thought of in a few main groups. One is cycle stock, basically the regular amount kept on hand to cover demand between two orders. There's also safety stock, which acts as a cushion when demand jumps unexpectedly or when deliveries take longer than planned. In some situations, companies prepare seasonal stock, building it up ahead of predictable peaks like holidays or seasonal sales. We also have pipeline stock, which simply refers to goods that have already been ordered but are still on their way. And then there's obsolete or idle stock, the leftover items that result from changes in demand or product updates, often a source of extra disposal costs.

Managing inventory also means dealing with a range of expenses. Among the most common are:

- Ordering or replenishment costs: administrative tasks, transportation, and setup costs linked to restocking.
- Stockout costs: the losses that come from running out of stock, such as missed sales, contract penalties, and dissatisfied customers.

The real challenge of inventory optimization is to cut the overall cost of holding and replenishing stock without sacrificing service levels. In complex and rapidly changing environments, though, relying solely on manual calculations or fixed models often proves inadequate and can even lead to poor results. This is where artificial intelligence becomes valuable, providing new tools to deal with demand volatility, product diversity, and the geographic dispersion of warehouses.

3.3.2 Traditional Inventory Management Models⁵³

EOQ Model

The EOQ (Economic Order Quantity) model is one of the most classic and well established tools in inventory management theory. First introduced by Ford W. Harris in 1913 and later refined by other scholars, it aims to determine the optimal order quantity each time stock is replenished, with the goal of minimizing the total cost of inventory management.

In this framework, total inventory cost is made up of two main components. The first is the ordering (or setup) cost, which decreases as the order size increases. The second is the holding cost, which instead rises with larger order quantities. The exact point at which the two curves intersect, where their sum is at its lowest, represents the Economic Order Quantity. The basic EOQ formula can be expressed as follows:

$$EOQ = \sqrt{\frac{2DS}{H}}$$

In the EOQ formula, D represents the annual demand, S the ordering cost for each order, and H the annual holding cost per unit. This formula makes it possible to calculate the quantity that, if ordered consistently, minimizes the total inventory cost in a system where demand is constant.

The model is based on a set of quite strict assumptions. It works under the idea that demand is perfectly known in advance, constant, and evenly spread throughout the year. Lead time is considered fixed, so no delays or stockouts ever occur. Both ordering and holding costs are assumed to stay the same over time, and the price per unit does not change — meaning no bulk

discounts are taken into account. Lastly, the model assumes that items never become obsolete or deteriorate while in storage.

Although the EOQ model is still a solid theoretical reference and remains useful for small and medium-sized businesses, its practical relevance can be limited in real-world settings where demand is variable, lead times are uncertain, and supply chains are more complex.

(s, S) Model – Minimum Stock Policy with Replenishment to Maximum Level

The (s, S) model is one of the main stochastic approaches to inventory management and is widely used in industrial practice. It defines a threshold-based replenishment policy: an order is placed only when the inventory level drops below a predefined threshold, called s or minimum stock level, and the quantity ordered is simply enough to bring the inventory back up to the maximum level, called S .

Put simply, when the current inventory level is lower than s , the quantity ordered is equal to S minus the current inventory. If the inventory level is equal to or higher than s , no order is placed. In this framework, the current inventory is represented by I , the reorder point by s , and the maximum stock level by S , which is often referred to as the order-up-to level.

Unlike the EOQ model, the (s, S) policy is built to handle situations where demand is not perfectly stable. Instead of following a fixed order quantity, it places an order only when stock drops below a certain point and adjusts the size of the order to restore the desired level. This approach gives managers more flexibility and tends to work better in real-life conditions, such as multi-product distribution, retail supply chains, or warehouses where demand can be irregular.

Finding the right values for s and S is one of the trickiest parts when putting the (s, S) policy into practice. These two numbers aren't fixed once and for all: they depend on many factors that influence each other, like how demand is distributed over time, the cost of placing an order, how long and how variable the lead time is, the cost of keeping items in stock, and even the expected cost of running out of stock or failing to meet service levels.

In the traditional approach, estimating s and S usually means building a stochastic optimization model, often written as a dynamic programming problem. The problem is that as you add more products, longer planning horizons, and additional warehouse locations, the calculations quickly become harder to manage. This growing complexity is one of the main reasons why the method is not always practical on a large scale.

A further difficulty is that the outcome depends a lot on the quality of the data used. When past data are incomplete or fluctuate too much, the values of s and S suggested by the model can be misleading, and the whole policy may end up performing poorly. In these cases, artificial intelligence can be very helpful. By using real-time information and predictive methods, it can constantly adjust these values and keep the inventory policy closer to the real behaviour of demand.

3.3.3 Metaheuristic Optimization in Inventory Management

When standard mathematical tools like linear programming or integer programming are not enough, optimization can become a problem: in this type of scenario metaheuristic can become a good alternative. They are algorithms designed to look for acceptable solutions without relying on an exact formula. Many of them are inspired by what happens in nature, such as the way species evolve, how animals act together, or how physical systems react to change.

Problems suited for metaheuristics, such as inventory control or logistics network planning, are usually very large and messy. There can be an enormous number of possible choices, variables of different kinds to account for, and a lot of constraints that influence one another. To make things worse, the function that measures how good a solution is is often irregular and non-linear, which makes traditional analytical methods almost impossible to use.

One of the main reasons metaheuristics work well is that they don't settle too early for a solution that is just "good enough." Instead, they keep exploring different areas of the problem while also improving the most promising results found so far.

In supply chain management, these algorithms have been used successfully to deal with problems where the objective functions are non-linear, data are uncertain, and the network has several levels. In these cases, classical approaches are either not practical or too computationally expensive. Among the many metaheuristics that have been proposed, Genetic Algorithms are some of the most common⁵⁴.

Genetic Algorithms

The principle of natural evolution inspired in 1970s John Holland in the creation Genetics Algorithms, used as a computational model to simulate the adaptive processes observed in living organisms.

In a Genetic Algorithm, the search step takes place over a population of candidate solutions, called individuals: each individual corresponds to one possible way of configuring the problem, and its characteristics are expressed through a sequence of values known as a chromosome. The most common encoding scheme is binary, but in engineering and management applications it is common to use real-valued or hybrid (binary + real) representations, which allow for a more flexible handling of mixed variables and operational constraints.

The genetic process unfolds across multiple generations, during which the individuals in the population are evaluated using an objective function that assigns them a fitness score essentially a measure of how well each solution satisfies the problem's criteria. The fitter an individual is, the more likely it is to be chosen to produce the next generation of solutions. This happens through genetic operators that mimic natural evolutionary processes. In other words, selection reproduces the idea of survival of the fittest, while crossover takes fragments from two parent solutions and recombines them to create new offspring. Mutation, on the other hand, introduces small random changes, which helps maintain genetic diversity and prevents the search process from getting stuck too early in a suboptimal solution.

When we look at the specific case of the (s, S) inventory policy, often called the "min-max" or "order-up-to" policy, each individual in the population can be represented by a chromosome with two real-valued genes: the value of s and the value of S , with the obvious constraint that s must be smaller than S . The aim is to minimize a cost function that typically accounts for several components at once. These include the cost of holding inventory over time, the cost of placing new orders, and the penalties associated with stockouts or missed sales.

The fitness function essentially measures how effective a given pair (s, S) is at keeping total costs under control. In practice, every candidate solution receives a score based on the total simulated cost over a given time horizon, taking into consideration demand, which may fluctuate or be uncertain. Genetic algorithms then explore many different combinations of reorder thresholds and target stock levels, gradually homing in on those configurations that yield the lowest overall cost⁵⁵.

Are Genetic Algorithms AI?

Genetic algorithms are sometimes treated as a separate computational tool, yet today they are usually placed within the wider domain of artificial intelligence, particularly in the area of intelligent optimization. Russell and Norvig (2021), in what is probably one of the most influential AI textbooks, classify GAs as “machine learning methods without explicit training.” Put another way, they don’t use a fixed training dataset. Instead, they gradually improve candidate solutions through a kind of evolutionary trial-and-error, letting better ones survive and combining them to search for even better results. Instead, they keep refining potential solutions step by step, following an iterative process that takes inspiration from how species adapt and evolve over time.

This way of looking at them makes clear why they fit so well in the AI world. Like other approaches in this field, they can explore very large and complex solution spaces on their own and adjust to changes as they go. This flexibility is particularly useful when the problem cannot be expressed through a precise mathematical model, or when such a model would be too rigid to capture the variability of a real system.

3.3.4 Deep Reinforcement Learning

One of the most recent artificial intelligence paradigms is Reinforcement Learning (RL): it is based on the principle of having an agent interact with a dynamic environment to determine how to arrive at an optimal solution. Unlike supervised learning, in this case the agent has no best practices to follow, but only basic rules which to adhere to: knowing the limits it must impose on itself, it interacts with the environment and, through a process of trial and error, receives feedback in the form of rewards or penalties, which guides it in its search for the optimal choice. This sequential decision-making process is formalized through Markov Decision Processes (MDP). An MDP is defined by five elements (S, A, P, R, γ), where:

- S are all the possible states that the ambient can have.
- A is the set of actions available to the agent.
- P is the state transition function, which defines the probability $P(s'|s,a)$ of transitioning from state s to state s' by performing action a.
- R is the reward function, which assigns a numerical value $R(s,a)$ to the state-action pair.
- γ is the discount factor, with $0 \leq \gamma \leq 1$, which determines the importance of future rewards relative to immediate ones.

The agent's goal is to learn a policy $\pi: S \rightarrow A$ that maximizes the expected value of the discounted sum of future rewards, known as the “return”:

$$G_t = \sum_{K=0}^{\infty} \gamma^K R(s_{t+K}, a_{t+K})$$

Where s_t and a_t represent the state and the action during time t⁵⁶.

To solve an MDP, the most widely used algorithm is Q-learning: this model-free method allows the agent to learn the optimal policy even without knowing the actual dynamics occurring in the environment. This algorithm estimates the value of the function $Q(s,a)$, which in turn represents the expected value of the cumulative reward obtainable by the agent performing action a in state s. The algorithm's limitation is that it iteratively updates the function, storing the Q values for each state-action pair in an explicit table: this makes the algorithm unsuitable for complex environments with very large or continuous state or action spaces.

Deep reinforcement learning

The development of Deep Reinforcement Learning (DRL) can be understood as a response to the computational limitations inherent in tabular reinforcement learning techniques, achieved by combining advances in deep learning with established reinforcement learning paradigms. Although algorithms like Q-learning have long been regarded as robust solutions for environments with relatively compact state–action spaces, their scalability is severely challenged when the number of possible states increases combinatorially. The difficulty is particularly pronounced in environments that exhibit continuity, multimodality, or high dimensionality, conditions under which any attempt to enumerate the state space becomes not merely inefficient but computationally prohibitive.

In such contexts, DRL appears to offer a compelling alternative by leveraging deep neural networks as function approximators. These architectures are able, at least in principle, to represent highly complex mappings and to process structured inputs without requiring the manual specification of salient state features. Rather than exhaustively enumerating every possible configuration, the network generalizes across similar states, allowing for more scalable estimation of value or policy functions.

This capacity is not merely a technical convenience but has important implications for domains such as supply chain management. Decision-making within such systems is conditioned by a dense network of interdependent factors—ranging from demand volatility and inventory positioning to lead-time variability, logistics costs, seasonal effects, and capacity limitations dispersed across multiple nodes of the network. Deploying DRL in this context may be interpreted as part of a wider research trajectory aimed at modelling decision processes in complex adaptive systems. Traditional analytic approaches, while often elegant in their formalism, tend to falter when confronted with the nonlinear interactions and stochastic fluctuations characteristic of real-world supply chains, which makes the case for learning-based, data-driven methods particularly compelling.

When you look at these points together, it becomes clear that DRL isn't just about faster computation. It seems to be changing how we even think about decision-making in supply chains. Instead of treating operations as a problem you solve once and then assume will stay the same, DRL keeps learning as the system changes, as markets swing up and down, as disruptions ripple through the network, as coordination between nodes shifts in unexpected ways. This kind of approach feels closer to a “living models” of supply chains: models that grow and adapt alongside the systems they represent. And that perspective challenges a lot of the old reliance on fixed parameters and static optimization, which have always struggled to keep pace with the real world⁵⁷.

3.4 Logistics

3.4.1 Introduction

Distribution logistics is the part of the supply chain that deals with moving goods along the distribution network. This can mean transporting raw materials to the manufacturer, sending semi-finished goods to other plants, or delivering finished products to distribution centers or directly to customers. It covers a range of activities such as transport management, route planning, delivery scheduling, and fleet coordination. The main priority is to guarantee the reliability of the distribution flow, and right after that its efficiency while keeping time, costs,

and environmental impact to a minimum, and staying within contractual limits and service expectations.

In today's fast-moving and interconnected global market, distribution logistics has become a strategic factor in business competitiveness. More and more companies are adopting advanced technologies to keep track of goods in real time, automate decision-making, and strengthen the resilience of their distribution networks. The way these systems perform has a clear effect on customer satisfaction, the sustainability of the entire supply chain, and on how quickly businesses can react when market conditions change.

3.4.2 Vehicle Routing Problem (VRP)

Introduction to VRP

The Vehicle Routing Problem (VRP) is one of the central challenges in operations research applied to logistics. It consists of finding an optimal set of routes for a fleet of vehicles that must serve a group of customers starting from one or more depots. Each customer must be visited exactly once, every vehicle has a limited capacity, and the usual goal is to minimize the overall transportation cost, whether measured in terms of distance travelled, time spent, or the number of vehicles used. The problem was first formally introduced by Dantzig and Ramser in 1959 to optimize fuel distribution, and since then it has become a reference model for a wide range of real-world applications: from urban freight distribution, to waste collection, to home healthcare services.

When looking at it from a computational perspective, the VRP falls into the class of NP-hard problems. In simple terms, the time it takes to find an exact solution grows extremely fast as the number of customers increases, to the point that solving large instances with exact algorithms becomes practically impossible. This is why most of the research in the field has moved toward heuristic and metaheuristic approaches and, more recently, toward techniques that use artificial intelligence and machine learning. The aim is to find solutions that are good enough and fast enough to be applied in real operational settings.

The VRP is important not only because of how often it appears in real-world logistics, but also because it has become a standard testbed for developing and evaluating new algorithms in combinatorial optimization⁵⁸.

The Two-Echelon Capacitated Vehicle Routing Problem: Definition

To discuss the "Two Echelon Capacitated Vehicle Routing Problem (VRP)", the chapter will refer to the teaching material written by Professor Gianpaolo Perboli, that is titled "Il problema di instradamento dei veicoli". This research was presented during course Decision Making and AI for Business Change at the Politecnico di Torino: it provides a clear and rigorous introduction to the problem and mirrors the approach used throughout the course.

This choice makes it possible to present the topic in a way that is consistent with the educational path followed and with the methodological framework acquired during the program⁵⁹.

The Two-Echelon Capacitated Vehicle Routing Problem (2E-CVRP) is a major extension of the classic VRP model, designed to better capture the needs of modern urban distribution and multi-echelon logistics systems. In this setting, goods are not delivered directly from the central depot to the final customers but move through a two-tier network. This network includes satellites, or intermediate platforms, that act as transit and consolidation points. First-level vehicles transport loads from the depot to the satellites, and from there, smaller vehicles handle the last-leg deliveries to customers. By structuring the network this way, it becomes easier to manage the movement of goods in busy urban areas, ease traffic in city centres, make deliveries

more punctual, and get the most out of the available resources. The overall objective is to minimize the total cost of the distribution system, which includes not only the distance travelled and the number of vehicles used, but also any penalties for violating constraints such as vehicle capacity or synchronization requirements.

The two echelons correspond to the two levels of the network: the first connects the main depot to the satellites, while the second links the satellites to the customers. It is assumed that all vehicles operating on the same level have identical transport capacity. However, the number of vehicles available at each satellite is not fixed in advance. Each customer has a fixed demand, denoted as d_i . It is also assumed that a customer's demand is less than or equal to the capacity of a second-level vehicle, and that this demand cannot be split across two routes within the same level.

The figure below shows an example of a 2E-CVRP network: each arc represents a possible route between two nodes, along with the corresponding notation.

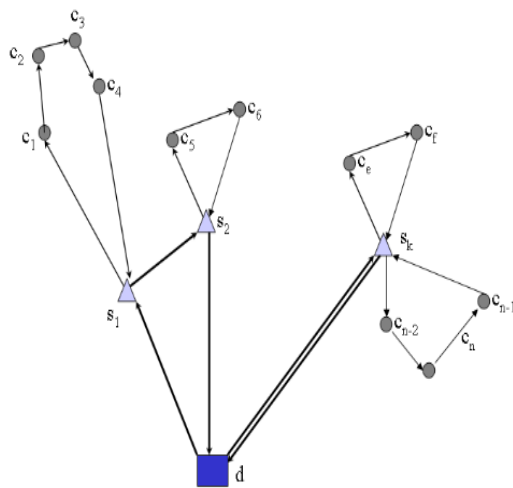


Figure 3.1 - 2E-CVRP network

Base model for 2E-CVRP

The goal of the model is to minimize the total cost of the distribution system. This includes the travel costs of first- and second-level vehicles and, optionally, the loading and unloading costs at the satellites. The decision variables cover several aspects of the problem: they include the route selection variables x_{ij} and $y_{i,j}^k$, the customer-to-satellite assignment variables z_{ik} , and the flow variables Q_{ij} , together with some auxiliary variables used to balance demand. The model also incorporates a comprehensive set of constraints. These ensure that vehicle and satellite capacities are respected, that each customer is assigned to exactly one satellite, and that the amount of goods delivered matches the total demand. They also guarantee the consistency of flows across incoming and outgoing arcs at intermediate nodes and eliminate subtours that do not include the depot or satellites.

By adding what is called "Valid Inequalities" the model improves. These are extra constraints that don't change which solutions are allowed but make the relaxed version of the model describe the problem more accurately. In practice, they make the model stronger and help branch-and-bound algorithms work faster. The authors divide these inequalities into two main types: edge cuts, which stop the creation of closed routes that aren't connected to the depot or

satellites, and flow constraints, which reinforce the capacity limits already set in the basic model. Adding these constraints in a controlled way leads to tighter lower bounds in the linear relaxation, reduces the number of fractional solutions, and speeds up computation.

This dual structure, a basic model complemented by valid cuts, is a well-established strategy in the combinatorial optimization literature and proves particularly effective for the 2E-CVRP, where the explosive complexity of real-world instances makes it essential to strike a balance between theoretical precision and computational tractability.

Math-Based Heuristics for 2E-CVRP

Perboli, Tadei, and Vigo then introduce two model-based heuristics for solving the Two-Echelon Capacitated Vehicle Routing Problem (2E-CVRP). These heuristics are developed starting from the continuous relaxation of the mathematical model presented in the previous chapters. The main idea behind both approaches is to use the information obtained from the relaxed version of the problem, where binary variables are treated as continuous, to effectively guide the construction of feasible integer solutions. In particular, the proposed heuristics focus on the variables z_{ik} , which define the assignment of each customer i to a satellite k . Once this assignment is fixed, the problem breaks down into independent subproblems, each of which can be solved as a traditional CVRP for its corresponding level of the network. The two strategies described are as follows:

1. Diving-Based Heuristic

This technique follows an iterative process that starts from the solution of the relaxed model. At each step, it checks the variables z_{ik} : those with values close to 1, which means they are very likely to be part of the final solution, are fixed to 1, while those with values closer to 0 are set to 0 and left out. The approach takes inspiration from the idea of pseudocosts, which roughly estimate how much each variable affects the optimal value of the objective function. After each round of fixing variables, the model is solved again in the hope of finding a feasible solution. If the model turns out to be infeasible, the algorithm “restarts” by changing which variables are fixed and then tries again. This heuristic is quite effective at exploring the solution space and does so with a reasonable computational cost.

2. Semicontinuous Heuristic

The second heuristic uses a simplified model called the semicontinuous 2E-CVRP: in this model some of the binary variables, such as $y_{i,j}^k$, are relaxed. The process works as follows. First, the continuous relaxation of the simplified model is solved, and the binary variables z_{ik} are fixed to integer values. Then, a MIP solver⁶⁰ is run on the model with a reduced subset of variables, within a preset time limit (for example, 60 seconds), saving the best solution found. For each candidate solution, the corresponding CVRP instances for the first and second levels are built and solved using a dedicated solver, again with a fixed time limit. Finally, the best solution from the entire process is selected. Although they do not guarantee optimality, the solutions produced by these two heuristics are competitive in both cost and computing time when compared to exact methods, making them especially suitable for real-world operational applications.

Overall computational results

In terms of solution quality, the tests carried out on instances with up to 50 customers and 4 satellites showed that the heuristics were able to find solutions with an average gap of less than 2% compared to the known optimum or the best available lower bound. In some cases, the gap dropped below 1%, which confirms that the approach works well even in more complex scenarios. As for computation times, both heuristics performed much faster than the exact

methods. The diving-based heuristic found good solutions in under 30 seconds in most cases: instead the semicontinuous heuristic took a bit longer but delivered more stable results in terms of average gap. Another important result is robustness: both heuristics gave consistent outcomes across multiple instances, with only small variations between different runs. Finally, when compared with exact methods, it becomes clear that solving the full model exactly quickly becomes computationally expensive, even for medium-sized instances. The heuristics, on the other hand, manage to deliver solutions of comparable quality in a fraction of the time, making them much more practical for real-world applications.

Why include heuristics in AI?

2E-CVRP and the related solution techniques fits perfectly in the broader field of artificial intelligence applications in for supply chain. Logistics flow optimization is, in fact, one of the main areas where AI is applied, and it covers not only machine learning methods but also advanced algorithmic approaches for automatically solving complex decision-making problems. The math-based heuristics analyzed in the context of the 2E-CVRP fall into this category. They are computational tools capable of exploring very large and diverse solution spaces, delivering effective solutions within reasonable computing times for real-world logistics problems, which are often full of constraints and discrete variables.

Among artificial intelligence approaches, these techniques belong to the field of algorithmic operations research, which is widely recognized as a core element of modern intelligent decision-support systems. What makes them so valuable is their ability to automate distribution planning, handle limited resources efficiently, and adapt to complex urban and multi-level logistics scenarios. This is why they have become an essential part of AI-based solutions currently used in the industry. For this reason, focusing on the 2E-CVRP is both methodologically sound and fully aligned with the purpose of this thesis: to explore how artificial intelligence can be applied to make supply chains more efficient, flexible, and responsive.

A concrete example of how the 2E-CVRP model can be advanced with AI techniques is provided by Yang (2024)⁶¹, who tackled the Two-Echelon Vehicle Routing Problem by combining Reinforcement Learning with heuristic methods. The model assumes a two-level network (main depot → satellites → customers) and introduces an agent that learns to make decisions for instance, which satellite to serve first or how to assign customers to local routes, based on a policy built through simulated experiences. The approach starts with an initial constructive procedure, then applies reinforcement learning to test and evaluate different assignment moves between customers and satellites. Step by step, the agent improves the solution, considering both travel costs and the efficiency of the overall logistics flow. The results show that this AI–heuristic integration can achieve better solutions than traditional 2E-VRP methods in benchmark cases, particularly by lowering total costs and handling customer-to-satellite assignments in a more dynamic way, while keeping computation times competitive. This development illustrates well how AI can push 2E-CVRP models beyond conventional heuristics: the agent is able to learn from examples, adapt to new scenarios, and improve its performance through experience.

3.5 Computer Vision for Quality Control

3.5.1 Introduction

Product quality is one of those things that can really make or break a company's competitiveness and it obviously matters a lot for customer satisfaction too. In many traditional production processes, quality control is still done manually or just by taking random samples.

This approach is slow, expensive, and, let's be honest, pretty easy to mess up because of human error.

Thanks to the growth of Artificial Intelligence and Computer Vision, the situation has really changed: today, it's possible to use automated systems that inspect products in real time. This allows plants to check huge volumes quickly and with a level of accuracy that would be impossible for humans to match. These AI systems usually combine convolutional neural networks (CNNs) with the image processing algorithms: this combination allows to spot defects, anomalies, or anything that doesn't meet the expected standards. What's nice is that they don't just detect problems more reliably than manual inspections but they also collect data that can be used to keep improving the process over time. This means plants can identify recurring issues and even trace back the root causes of defects instead of than just reacting when they show up.

Integrating this kind of technology into the supply chain helps reduce waste, cut product recall costs, and makes the whole production process more resilient. Also, since everything happens in real time, plants can step in immediately when something goes wrong, stopping defects from spreading further down the line and keeping the quality level high.

3.5.2 Machine Vision System's components

For this paragraph and for the whole section on AI and Computer Vision-based Quality Control (3.5) the discussion will be based entirely on the work of Ettalibi, Elouadi, and Mansour (2024)⁶²: the text provides a comprehensive review of the main computer vision technologies and their industrial applications.

A Machine Vision System (MVS) is a technological framework that makes possible to automate visual inspection into an industrial production. Its architecture usually consists of three key components: lighting, image acquisition, and image/signal processing.

Lighting

This step is essential to get some clear and high-quality images, without noise. The goal is to highlight the features of the object being inspected: at the same time it has to reduce shadows and reflections as much as possible. Different setups can be used, such as front-lighting (direct illumination on the object, ideal for analysing surface details) or back-lighting (illumination from behind, useful for emphasizing shapes and edges). Depending on the application, besides traditional light sources, lasers, infrared lights, fluorescent lamps, or even X-rays can be used to inspect complex materials or elements that are not visible to the human eye.

Image Acquisition

Images are captured using sensors and cameras: they are often based on a Charged Coupled Device (CCD) or on a CMOS technology, which convert light into digital signals. Modern digital cameras offer high resolution and low noise making it possible to inspect very small and detailed objects and with even complex form. Choosing the right optics is crucial to find the right balance between field of view and precision, and it can also include optical filters or polarizers to improve image quality in challenging production environments.

Processing and Communication

Once the images are captured they are immediately sent to a processor: there, are performed image processing tasks such as filtering, segmentation, and edge detection. In more advanced systems, AI and ML models are also used to recognize and classify defects. Thanks to edge computing techniques and hardware accelerators (like GPUs, FPGAs, or Intel Neural Compute Sticks), this analysis happens in real time: this allows the immediate detection of non-

conformities and also allows to trigger automatic countermeasures directly on the production line. The results are then sent to supervision and traceability systems, enabling fast, integrated responses within the production flow.

3.5.3 AI Techniques for Quality Control

Artificial Intelligence techniques used in quality control aim to perform important tasks: in real time they have to automatically detect defects, anomalies, or non-conformities within production processes. In this context, AI is not just about classifying or predicting events, but it needs to work at high speed, with great accuracy, and under changing environmental conditions.

A key role is played by Convolutional Neural Networks (CNNs), which are widely used for visual recognition of surface defects and microcracks. In quality control, these networks are trained on datasets made up of images: they contain both good and defective products, so they can learn visual patterns that allow them to tell the difference. This approach is especially effective in situations where the difference between a “good” product and a defective one is very subtle and hard to catch even for the human eye.

Other supervised machine learning techniques are also used, besides CNN. Random Forests and Support Vector Machines (SVMs) are particularly suitable for classifying specific and recurring defects. A common example is the detection of colour defects, where images are pre-processed to normalize brightness and colour features are extracted in RGB or HSV spaces. These algorithms are valued for their fast inference times and their relatively easy interpretability: these features make them ideal for high-speed production lines and for situations where decision traceability is required (for instance, in the food and pharmaceutical industries). While they may be less powerful than deep networks, they offer a good balance between accuracy, computational cost, and ease of implementation in industrial environments.

Finally, combining AI with edge computing is a key step forward in quality control. Edge computing means processing data close to where it is generated: this is translated into processing data directly on the machine or at the network edge, instead of sending it to a central server or the cloud. This approach reduces the amount of data that needs to be transmitted, lowers processing latency and improves reliability: this is very important because the system can keep working even without a stable internet connection. This also means that AI algorithms can run directly on edge devices connected to the production line cameras. As a result, the system can make decisions in real time and react right away, for example by removing a defective part, notifying operators, or adjusting process settings. This makes quality control more reactive and adaptive, able to keep up with fast production rates and improve the overall stability of the line.

3.5.4 Colour Recognition and Measurement

Colour is one of the first quality indicators noticed by the consumer. A product that is too dark, too light, or with shades can be seen as defective, even if all its functional properties are perfectly fine. For this reason, color recognition and measurement play a key role in quality control across many industries, especially in food, cosmetics, plastics, and textiles. Traditionally, colour measurement was done through visual inspections by trained operators or with dedicated instruments like colorimeters and spectrometers. These methods still nowadays accurate but they are also slow, expensive, and hard to integrate into fast production lines. The introduction of computer vision systems has made it possible to measure colour automatically and in real time by analysing digital images captured along the production line.

Colour is usually encoded in the RGB (Red-Green-Blue) space: there each pixel is described by a combination of three numerical values representing the intensity of the primary components. To perform more advance applications they are used more robust colour spaces: HSV (Hue, Saturation, Value) or CIELAB. These separate the colour information from the illumination, allowing for more stable detection even when ambient lighting changes. After acquisition, images are calibrated using standard references (like a colour checker) to ensure that measurements are consistent and comparable over time.

From an algorithmic point of view, the detection of colour defects usually follows three main steps. The first one is the pre processing phase: here brightness is normalized and white balance is adjusted. Then the image goes through the segmentation phase which isolates relevant areas like the surface of the product. The last phase is the classification phase: there the colour values are compared with the acceptable ranges to spot any significant deviations. In this process is it possible to detect issues like discoloration, stains, contamination, or shade variations, while reducing false positives and ensuring that only truly defective items are discarded.

3.5.5 The importance of AI in Quality Control

In the context of Quality 4.0 automatic colour measurement helps in many different ways. It keeps the product standardized and it make sure that different batches and plants stay consistent. It also lowers scrap and rework costs by finding defects earlier during the process. On top of that, it improves traceability because color data can be saved and used later to make upstream processes better. Finally, it makes customers happier, since the product looks the same every time and matches what they expect. By combining visual recognition, AI, and colorimetry techniques, quality control becomes much more than a simple manual check. It turns into an integrated, automated, and even predictive system that can support the whole supply chain and make sure products meet the right standards from the very first step of production.

3.6 Summary and Outlook

This chapter gave an overall look at how artificial intelligence is being used in the supply chain: it has also shown why it has become so important for dealing with volatility, complexity, and strong competitive pressure. It started with demand forecasting, pointing out the limits of traditional statistical models and showing how machine learning and deep learning can add value by capturing non-linear patterns and adapting quickly when things change. For inventory management, it explained how metaheuristics and evolutionary algorithms, like Genetic Algorithms, can help optimize stock levels in complex networks with uncertain demand, offering a smarter alternative to classic deterministic formulas.

The section on logistics showed how AI has become a key tool for solving complex combinatorial problems like the Vehicle Routing Problem and its multi-level variants. Using mathematical models enhanced with smart heuristics makes it possible to get high-quality solutions in a time frame that works for real operations, enabling dynamic resource allocation and significantly improving both costs and service levels. At the same time, the use of computer vision and deep learning for quality control highlighted a major shift from manual, reactive inspection systems to automatic, predictive, and adaptive monitoring systems. These new approaches help reduce waste and make production processes more reliable.

Looking ahead, supply chains are expected to get smarter, more connected, and almost self-managing. Technologies like digital twins will let companies build virtual versions of their entire logistics network. This way, they can run simulations and test “what-if” scenarios before making big decisions. Generative models and reinforcement learning will allow strategies that

can adapt in real time. This means the supply chain will be able to adjust on its own when there are problems like supply shortages, sudden demand changes, or new regulations. With edge AI and distributed systems, the “brain” of the supply chain will move closer to where decisions are needed. This makes it possible to react much faster without waiting for everything to be processed by a central system. But this future won’t be without challenges. It will be necessary to make sure the data used in these models is high-quality, reliable, and secure. There’s also the issue of transparency understanding how algorithmic decisions are made and the need to train people so they have the right analytical and digital skills. Ethical AI governance, privacy protection, and following regulations will be essential to make adoption sustainable.

In the end, AI is a big chance to transform supply chains from reactive systems into predictive, adaptive, and resilient ecosystems. This can create value not only for the business but also for society and the environment over the long term.

CHAPTER 4 – Applied Case Studies: AI in Supply Chain Management

4.1 Introduction

In the previous chapters, the key ideas behind Artificial Intelligence and its main technologies were explored, along with the structure and challenges of today's supply chains. This fourth chapter moves one step forward by looking at real examples of how AI has been used in large companies. The goal is to show the practical benefits it can bring but also the problems that emerged, and the strategies that were used to overcome them.

By looking at selected case studies, it will be shown how AI has been integrated into 4 key areas: demand forecasting, inventory management, warehouse automation, and logistics optimization. These examples clearly show how technological innovation is changing the way companies manage their supply chains.

The objective of the chapter is twofold: on the one hand, to highlight the transformative potential of AI, and on the other hand, to shed light on the challenges linked to its large-scale implementation. The case studies will be analyzed with a critical approach, considering both the achieved results and the limitations and obstacles encountered, in order to offer a realistic and comprehensive picture of the impact of Artificial Intelligence on supply chains.

4.2 Walmart sales forecasting

4.2.1 Case study introduction

Today, predicting demand is very important to help companies work more efficiently and avoid wasting products. Walmart, one of the biggest retailers in the world, uses machine learning models to guess its future sales. This helps the company organize its inventory better and keep the right amount of stock.

An article written by Cyril Neba, Shu F. B. Gerard, Gillian Nsuh, Philip Amouda, Adrian Neba, F. Webnda, Victory Ikpe, Adeyinka Orelaja and Nabintou Anissia Sylla in 2024, in the *Asian Journal of Probability and Statistics*⁶³, examined the effectiveness of several machine learning algorithms applied to Walmart's historical sales data: these were collected over more than two years from many stores across the United States. The dataset included variables such as weekly sales, temperature, fuel prices, the Consumer Price Index (CPI), unemployment rates, and holiday flags. While the study reports significant results in the field of sales forecasting through machine learning, some limitations should be considered. First, the dataset covers a relatively short time period (2010–2012), which may reduce the validity of the predictions when compared to current market dynamics, which have changed significantly after global events such as the COVID-19 pandemic. Another limit is that the data comes only from some areas of the United States, so it might not show how people buy in other regions. Also, there are no socio-demographic data, so it is not possible to see if the predictions are fair for different groups of people. The study only considers time-based factors, like holidays.

Still, even with these limits, the study is a good base for creating more advanced prediction models for the retail sector.

4.2.2 Data exploration

Before starting the predictive modelling, I did an Exploratory Data Analysis (EDA) to see how the main variables were distributed and to check if there were any strange values. The analysis

focused on finding outliers in three important variables. This step is useful because unusual values can affect the accuracy of the models.

In retail sales, extreme numbers can happen because of promotions, holidays, mistakes in the data, or special situations. If these outliers are not managed, they can change how the model learns and give results that do not work well in general. For this reason, finding and, if needed, fixing outliers is an important step to build models that can give reliable forecasts even when the data changes a lot.

From the analysis of the dataset, it was found that weekly sales have an average of about 1.05 million units, but with a high standard deviation, showing that demand is quite unstable. Other factors like fuel price, CPI and unemployment rate change a bit over time, and these changes still matter, especially when there are economic or seasonal events. The link between these variables and sales is not very strong if we just look at simple linear correlations. However, the analysis showed that holidays clearly affect how people buy. This suggests that simple prediction methods may not be enough to really understand the data, so using more advanced machine learning models can help to catch more complex patterns. For this reason, more advanced machine learning models are needed, as they can find non-linear patterns and interactions between variables.

4.2.3 Modelling methodology

To make reliable sales forecasts, the study followed a clear and careful process, starting with data preprocessing. After the exploratory analysis, winsorization was applied to reduce the effect of extreme values that could make the predictions less accurate.

In simple terms, winsorization replaces the most extreme numbers with the nearest value inside a chosen limit, so the general shape of the data stays the same and no observations are removed. This step was done on the most important variables, like Weekly Sales, Holiday Flag, Temperature, Fuel Price, CPI, and Unemployment, to make the data more stable and easier to use for building predictive models.

To make reliable sales forecasts, the study followed a clear and careful process, starting with data preprocessing. After the exploratory analysis, winsorization was applied to reduce the effect of extreme values that could make the predictions less accurate. In simple terms, winsorization replaces the most extreme numbers with the nearest value inside a chosen limit, so the general shape of the data stays the same and no observations are removed. This step was done on the most important variables, like Weekly Sales, Holiday Flag, Temperature, Fuel Price, CPI, and Unemployment, to make the data more stable and easier to use for building predictive models.

For the modelling part, different machine learning algorithms were tried. The process started with simple methods like Linear Regression, Multiple Regression, and Generalized Linear Models, and then moved to more advanced ones. A Decision Tree was also tried, since it gives an easy-to-interpret model, as well as ensemble methods like Random Forest and Gradient Boosting Machine, which can capture more complex relationships between variables.

Finally, optimized boosting methods like XGBoost and LightGBM were also tested. These models are fast, usually very accurate, and can deal with missing data while helping to avoid overfitting. All the models were trained with their default parameters, which made sense given that the dataset was not very big and it was important to keep a good balance between model complexity and generalization.

4.2.4 Models results

To understand which model was better, the study used 3 indicators. The first one was Mean Absolute Error (MAE), that measures the average absolute difference between the predicted and the real values. The second one was Root Mean Squared Error (RMSE), which gives more weight to big errors and is more sensitive to extreme deviations. The last one was R^2 coefficient: it shows how much of the data variability is explained by the model; values close to 1 mean the model predicts very well. The results [Table 4.1] showed clearly that the more advanced machine learning models performed better than the traditional linear ones. Linear Regression, Multiple Regression, and GLM gave very similar results, with a MAE of 35,632.510 and an RMSE of 80,153.858. These values are too high to describe sales in a good way. The Decision Tree was the worst, with a MAE of 77,388.082 and an RMSE of 93,721.066: that showed that this type of model struggles to catch complex relationships. On the other hand, ensemble models like Random Forest and Gradient Boosting Machine were much more accurate. Random Forest reached a MAE of 12,238.782 and an RMSE of 19,814.965, while GBM achieved a MAE of 10,839.822 and an RMSE of 14,110.831. The best results came from the optimized boosting models. XGBoost had a very low MAE of 1,226.471, an RMSE of 1,700.981, and an almost perfect R^2 (0.9999900). LightGBM also performed very well, with a MAE of 1,692.640 and an RMSE of 2,297.930.

These results confirm that boosting models can find non-linear patterns in the data and reduce prediction error much better than simpler models.

Model	MAE	RMSE	R Squared
Linear Regression	35632.510	80153.858	0.9760562
Multiple Regression	35632.510	80153.858	0.9760562
GLM	35632.510	80153.858	0.9760562
Decision Tree	77388.082	93721.066	0.9696449
Random Forest	12238.782	19814.965	0.9986431
GBM	10839.822	1700.981	0.9993119
XGBoost	1226.471	1700.981	0.9999900
LGB	1692.640	2297.930	0.9999818

Table 4.1 – Results of predictive modeling

The study also looked at bias [Table 4.2], meaning the systematic tendency of a model to predict values that are too high or too low compared to the real ones. The linear models showed a positive bias of about 1,211.869, while the Decision Tree had an even higher bias (1,691.079). Random Forest showed a negative bias: that means it constantly underestimated sales. GBM, XGBoost, and LightGBM instead had bias values very close to zero, with XGBoost showing the lowest value (−7.548432): it means its predictions were very well balanced.

Model	MAE
Linear Regression	1211.869
Multiple Regression	1211.869
GLM	1211.869
Decision Tree	1691.079
Random Forest	-965.4004
GBM	-95.95693
XGBoost	-7.548432
LGB	99.07631

Table 4.2 – Model bias

Finally, a fairness analysis was done to check if the predictions were consistent across specific subgroups of the dataset, in this case holiday weeks and non-holiday weeks. The results showed [Table 4.3] that linear models gave comparable predictions in the two subgroups, while the Decision Tree was more accurate in non-holiday periods. The more advanced models were still the most accurate overall, but they had lower MAE values for non-holiday weeks compared to holiday weeks. This shows that the models struggled more to catch the changes that happen during holidays. In practice, it would be useful to improve them so they can predict demand peaks in those periods more accurately.

Model	Subgroup	MAE
Linear Regression	Holidays	58198.315
	Non-Holidays	24607.672
Multiple Regression	Holidays	33915.546
	Non-Holidays	11297.671
GLM	Holidays	58198.315
	Non-Holidays	12091.195
Decision Tree	Holidays	33915.546
	Non-Holidays	10744.609
Random Forest	Holidays	58198.315
	Non-Holidays	1146.513
GBM	Holidays	33915.546
	Non-Holidays	1232.555
XGBoost	Holidays	84273.385
	Non-Holidays	1757.721
LGB	Holidays	76864.200
	Non-Holidays	1687.688

Table 4.3 – Model Fairness

4.2.5 Managerial implications, limitations and future prospects

The results of this study are very useful for a retail company like Walmart. When sales forecasts are more accurate, it is easier to manage stock: this reduces the risk of empty shelves and also the problem of having too many products in the warehouse. This means to lower storage costs and make customers happier. Simple models like regressions and decision trees are helpful because managers can understand which factors affect demand and plan strategies for different products or times of the year. Ensemble methods like Random Forest and GBM make the predictions more accurate and show which variables are the most important, giving managers good information to make decisions. Instead, advanced models like XGBoost and LightGBM give very accurate results and tools like SHAP values can explain them. These models can help managers plan promotions and marketing campaigns, use resources at the right time, and take full advantage of the demand during holidays.

However, this study has some limits. The dataset is small and only covers 2010 to 2012. Because of this problem, it doesn't reflect today's market, which changed a lot after events like COVID-19. The data is also collected only from some parts of the United States and does not include information about people: so the study cannot check if the models work well for different groups. The data is not equally detailed for all stores, and this could make the results less accurate and a bit noisy.

In the future, the dataset should include more recent data and more variables, so the models can follow how the market is changing. Another good idea could be to try deep learning methods, like recurrent neural networks, to better catch time trends and complex patterns. It is also important to keep checking bias and fairness, especially during holiday weeks, to be sure the models stay fair even when demand is very different from normal. Doing this can help retailers use predictive analytics in a smarter way and make faster, data-based decisions in a very competitive market.

4.3 Reinforcement Learning for inventory management

4.3.1 Introduction

In this chapter will be shown how Reinforcement Learning is used for inventory management, with two points of view: the first one is an academic from research and one more practical. The study by Sultana et al. (2020)⁶⁴ shows how multi-agent reinforcement learning can be used to manage restocking in networks with many products and nodes. What is special about this study is that both the problem and the solution come from a real company case: the companies can't be named for privacy reason. This makes the work very close to real life and useful for getting ideas that can be applied in practice. The results show that it is possible to move from only predicting demand to using a decision system that reduces stockouts and waste and respects capacity limits.

The article “Real-World Success Stories: How Top Companies Are Optimizing Inventory with AI in 2025”⁶⁵ gives an overview of the most recent uses of AI in inventory management: it shows numbers and success stories from companies like Amazon, Walmart, Zara, and Toyota.

Since companies rarely share the technical details of their AI systems, because of competition reasons or because the solutions are often very complex, combining an academic study based on a real case with an article that reports real results is a good approach. This way, it is possible to get a complete view that mixes theory and practice and gives a solid base to discuss the impact of AI on inventory optimization.

4.3.2 Reinforcement Learning for Multi-Product and Multi-Node Inventory Management in Supply Chains

Problem introduction

Inventory management with multi products and several locations is a very complex problem. It needs coordination of reorder decisions between different levels of the network and at the same time must respect capacity limits, logistic costs, and deal with demand uncertainty. The study by Sultana, Gosavi and Das (2020) focuses on this problem and shows how reinforcement learning can be used in a system where one central warehouse supplies three stores. The problem and the solution were adapted from a real business case, which makes the study practical and close to real situations.

In the proposed model, each store has a limited storage capacity, and the central warehouse also has a total capacity limit. The transport from the warehouse to the stores is restricted by volume. The demand for each product is random and changes over time, and the replenishment times are different for the various products. This creates a risk of both stockouts and having too much inventory.

The problem becomes even harder because there are hundreds of products to manage at the same time, and it is necessary to coordinate decisions between the store level and the warehouse level. The main goal of the study is to find reorder policies that maximize the number of satisfied sales and reduce waste, especially for perishable products. At the same time, the solution must respect storage limits and keep a fair distribution of resources between the different product categories. This approach tries to reproduce the real challenges of distribution networks, where it is always necessary to find a balance between product availability, logistics efficiency, and cost reduction.

Modelling and formulation RL

The system is modelled as a Markov decision process where, for each store $j=1,\dots,S$ and product $i=1,\dots,P$, the inventory after reordering is given by

$$x_j(t)^+ = x_j(t)^- + u_j(t),$$

where $x_j(t) \in [0,1]^P$ is the vector of normalized stocks and $u_j(t) \in [0,1]^P$ The required warehouse replenishment quantities. The store-level constraints are:

$$0 \leq x_j(t) \leq 1, \quad 0 \leq u_j(t) \leq 1, \quad 0 \leq x_j(t)^- + u_j(t) \leq 1, \quad v^T u_j(t) \leq v_{max,j}$$

v is the vector of unitary volumes per product and $v_{max,j}$ to transport capacity to the store j . Sales achieved $w_t(t+1)$ update the stock at the end of the period according to

$$x_j(t+1)^- = \max\{0, x_j(t)^+ - w_t(t+1)\},$$

while a predictor provides the estimate $\hat{w}_j(t+1) = f(w_t(0:t))$.

The store's local objective includes multi-objective components: out-of-stock penalties, waste (for perishables), and fairness between products. The cost for store j at time t is

$$C_{j,st}(t) = \frac{p_{j,empty}(t)}{p} + \frac{1}{p} \sum_i q_{waste,ij}(t) + \Delta x_j(t)^{.95-.05},$$

Where $p_{j,empty}(t)$ it is the number of products with zero stock, $q_{waste,ij}(t)$ the amount wasted and $\Delta x_j(t)^{.95-.05}$ the percentile differential $95^\circ-5^\circ$ on stocks to measure the equity between references. At the warehouse level, the availability $\chi(t) \in [0,1]^P$ imposes a budget constraint on shipments:

$$\sum_{j=1}^S a_j u_j(t) \leq \chi(t),$$

with a_j normalization factors. The warehouse is replenished every n periods through action $\mu(nt)$ (one-day lead time), with dynamics

$$\chi(n(t+1))^+ = \chi(n(t+1))^- + \mu(nt)$$

And constraints $0 \leq \chi \leq 1, \quad 0 \leq \mu \leq 1, \quad 0 \leq \chi^- + \mu \leq 1$.

A crucial element of the model is the reward function, which is the signal on which the learning algorithm is based. The global reward is constructed as a linear combination of warehouse costs and point-of-sale costs:

$$G_t = \sum_{k=0}^{\infty} \gamma^k [g_1 C_{wh}(t+k) + g_2 C_{st}(t+k)], \quad C_{st} = \sum_{j=1}^S C_{j,st}$$

This objective, which the algorithm seeks to maximize, incentivizes policies that reduce costs related to waste, stockouts, and excessive capacity usage, while maintaining a sufficient level of inventory to meet demand. Policy convergence, observed in the training curves, corresponds to reaching an average reward value that no longer improves significantly, indicating that the agents have learned stable stationary behaviour.

In this framework, the problem is solved with multi-agent reinforcement learning: store agents and the warehouse agent learn local policies $(x_j, \hat{w}_j, \chi) \mapsto u_j \in (X, \chi, \hat{\chi}, \hat{W}) \mapsto \mu$ respectively, where X collects the stocks by store-product and $\hat{\chi}, \hat{W}$ These are state forecasts and aggregate demand.

For the learning part, Advantage Actor-Critic (A2C) is adopted. At the store level, actions are quantized to ensure implementability and compliance with constraints: the actor selects from 14 discrete levels of normalized reordering.

{0, 0.005, 0.01, 0.0125, 0.0175, 0.02, 0.03, 0.04, 0.08, 0.12, 0.2, 0.5, 1},

While at the warehouse level, the action is binary (order/not order each product). Choosing quantization reduces the complexity of the action space, improves training stability, and keeps decisions interpretable for the business user.

Experimental setup and tested scenarios

The authors tested their approach with a set of simulation experiments using the public Instacart Market Basket Analysis dataset, which has about 3 million orders and more than 49,000 products. From this dataset, they selected the most frequent products and created three scenarios with increasing complexity: the first with 50 products, the second with 220, and the third with 1,000 products. All three scenarios have three stores supplied by one central warehouse. For each scenario, capacity limits were set for the warehouse and the stores, as well as transport limits and different replenishment times for each product. This was done to make the simulation as close as possible to a real system.

For each scenario, the multi-agent reinforcement learning model was first trained and then tested on a separate set of data that was not used before. The main reason was to check if the learned policies could generalize to new data well. The training used the Advantage Actor-Critic (A2C) algorithm, with network weight sharing between agents. This made the model scalable for a large number of products and stores.

The system performance was measured using several metrics: the number of stockouts per period, the amount of wasted product, the respect of capacity limits, and a fairness measure. This fairness was calculated as the percentile difference (95th–5th) of the stock levels, to make sure the inventory was distributed fairly between products.

Finally, the results from reinforcement learning were compared with two reference strategies: the s-policy, a reorder rule based on fixed thresholds, and a clairvoyant policy, which is an ideal benchmark that assumes perfect knowledge of future demand.

This comparison makes it possible to measure the improvement over traditional methods and to see how far the results are from a theoretically optimal solution.

Results and implications

The first key result is shown in Figure 4.1 and is about the training strategy. Training the store agents first, independently, and then training the warehouse agent leads to better global policies than training them all together or using shared rewards: this is shown by the line reaching a higher value of mean reward in an also shorter time. When the store agents are trained together with the warehouse, they tend to “ask for less” to adapt to the low stock in the warehouse. This reduces sales and makes the total reward worse. With sequential training, the stores can learn more aggressive policies that focus on meeting demand, and then the warehouse is trained to support this demand efficiently. This improves the whole system. This result is very important for real-world use: in a real deployment, it is better to train and validate the store reorder policies first, to make sure customer demand is satisfied, and only after that optimize the warehouse level to reduce costs without hurting product availability.

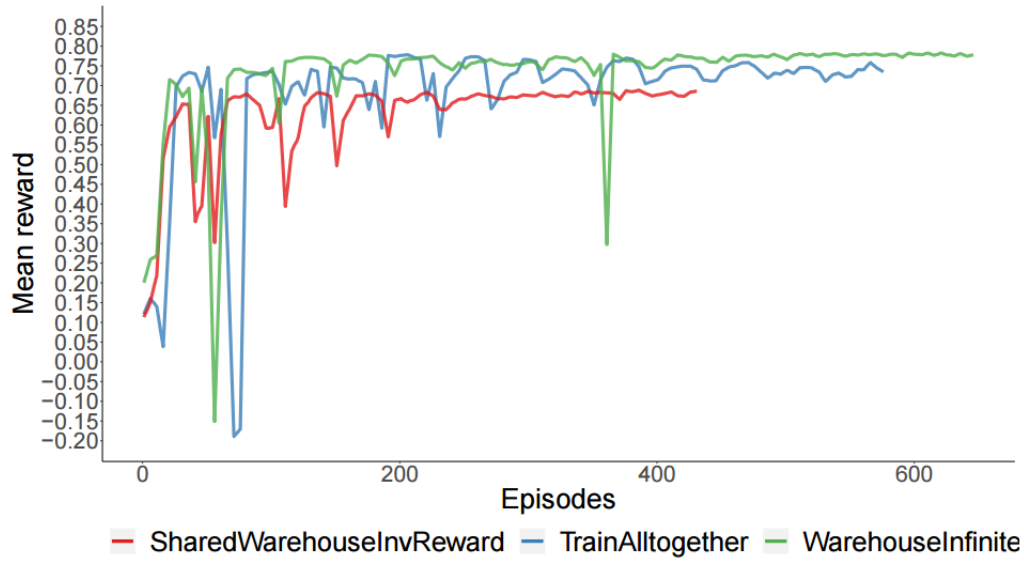


Figure 4.1 – Training strategy

The test results, shown in Table 4.4, confirm that the model can generalize well. In all three scenarios (50, 220, and 1,000 products), the multi-agent RL (MARL) performs better than the constant s-policy, both in terms of total reward and in reducing stockouts and waste. In some cases, it even score a higher results than the clairvoyant policy, which has perfect knowledge of future demand. This is an important result because it shows that the algorithm not only deals with demand uncertainty but also uses the network dynamics to improve how stock is allocated.

Methods	Warehouse	Store1	Store2	Store3
50 products				
RL	0.451	0.793	0.811	0.805
Heuristic	0.392	0.617	0.671	0.690
Clairvoyant	0.552	0.766	0.838	0.817
220 products				
RL	0.703	0.783	0.801	0.784
Heuristic	0.738	0.610	0.695	0.692
Clairvoyant	0.769	0.716	0.839	0.836
1000 products				
RL	0.702	0.796	0.820	0.818
Heuristic	0.660	0.591	0.679	0.669
Clairvoyant	0.735	0.750	0.833	0.680

Table 4.4 – Results on testing data

The component analysis, shown in Figures 4.2 and 4.3, gives a more detailed view of how the system behaves. For the stores [Figure 4.2] the reward linked to reducing stockouts increases during training and reach a level comparable to Clairvoyant policy. Instead the average inventory level stabilizes at moderate values, very close from the Heuristic value. This helps avoid overstocking and waste in the store. For the warehouse, the reward for rejected orders [Figure 4.3] increase step by step, showing that the policy has learned to prevent saturation and keep a steady flow of replenishment.

Overall, the system reaches a balance between three goals: reducing stockouts, limiting waste, and respecting capacity limits, while maximizing the total reward.

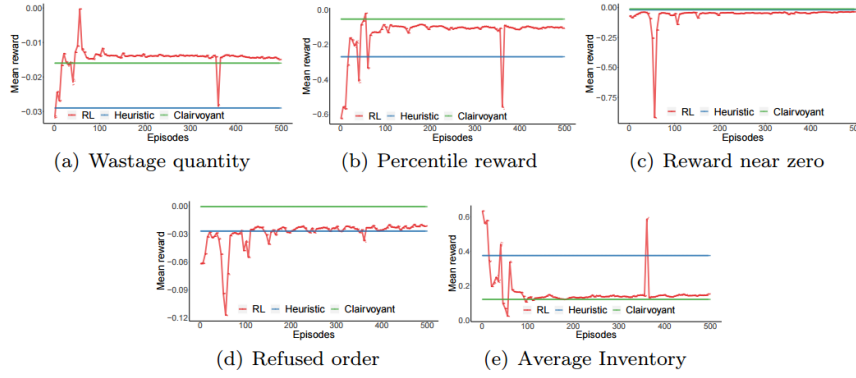


Figure 4.2 - Individual components of store replenishment rewards for Store 1 with 220 products

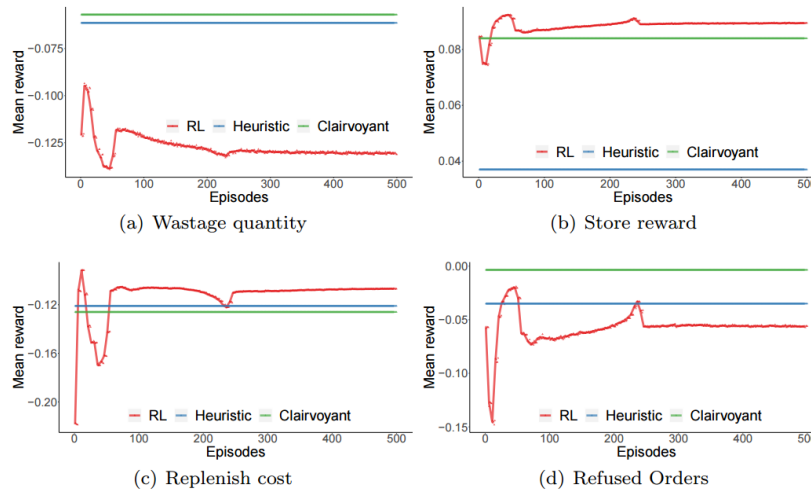


Figure 4.3 - Individual components of warehouse rewards, for 220 products, 3 stores.

From the point of view of adaptability, the transfer learning results (shown in Table 4.5) show that the learned policies can generalize and are not tied to the exact number of products or nodes used in training. In practice, a model trained on 50 products was applied directly to a scenario with 70 products, without any extra training, and reached almost the same performance as a model trained from scratch on the new scenario. This means that the agent learned a reorder strategy based on general state features (like stock level, expected demand, and remaining capacity) and not on fixed details of the product assortment.

Methods	Warehouse	Store1	Store2	Store3
70 products				
RL (transfer from 50 prod)	0.428	0.764	0.811	0.796
RL (trained on 70 prod)	0.431	0.784	0.819	0.812
Heuristic	0.299	0.554	0.672	0.673
Clairvoyant	0.498	0.779	0.837	0.825

Table 4.5 - Transfer learning of warehouse and stores on additional products

The same result was seen when a new store was added [Table 4.6]. The policy trained on three stores was able to handle a fourth store that was never seen during training, without changing or re-optimizing the neural network weights. The existing agents simply adjusted their decisions to the new demand distribution: it means that the multi-agent approach is both robust and modular.

Methods	Warehouse	Store1	Store2	Store3	Store4
50 products					
RL	0.446	0.792	0.803	0.806	0.787
Heuristic	0.373	0.619	0.672	0.693	0.651
Clairvoyant	0.544	0.767	0.838	0.818	0.750

Table 4.6 - Transfer learning of warehouse on an unseen Store 4, with 1.75x capacity of Store 1

This feature is very important for real-world deployment. It helps cut the cost of model maintenance and retraining, which would normally be needed every time the product assortment or the network setup changes (for example new products, store openings or closings, changes in capacity). In dynamic environments like e-commerce and omnichannel retail, where these changes happen often, the ability to transfer already trained policies to new scenarios keeps operations running smoothly and reduces the risk of performance getting even worse.

Implications

From a theoretical point of view, this study gives an important contribution to the literature on hierarchical control and distributed optimization. It shows that an approach based on Advantage Actor-Critic with action quantization can scale to hundreds or even thousands of products while staying stable and converging. From a practical and managerial perspective, the results suggest that RL systems like this can be used in real situations to reduce waste, increase product availability, and respect capacity limits without constant manual work. The fact that the policies can be reused for new network configurations, together with the suggestion to train stores first and then the warehouse, makes this method cost-effective and easy to scale. This can give a competitive advantage to retail and logistics companies working in markets with changing demand.

4.3.4 From Research to Practice: Real-World Evidence

These results match what is happening in recent industrial experiences reported in the article “Real-World Success Stories: How Top Companies Are Optimizing Inventory with AI in 2025”. According to the report, more than 80% of companies see inventory management as one of their main operational challenges. The adoption of AI-based systems has led to inventory cost reductions of 10–15% and supply chain efficiency gains of 20–25%.

The case studies in the article show concrete examples. Amazon has implemented a predictive system based on machine learning that analyzes sales, supplier lead times, and market trends in real time. This led to 35% fewer stockouts, lower transport costs, and a 20–25% increase in customer satisfaction. Sales also went up by about 5–7%.

Walmart uses autonomous robots to check shelves and update stock levels in real time. This cut counting errors, reduced stockouts by around 15% and overstocking by 20%, and brought a 3% increase in sales together with better customer loyalty.

Zara has focused on an AI forecasting system for fast fashion, which uses data from social media, weather, and market trends. With this system, Zara was able to cut extra inventory by around 40% and reduce unsold items by 25%. This also made their process more sustainable.

Toyota brought AI and IoT into its just-in-time system. Using predictive algorithms, it can now find supply problems earlier and keep stock levels under control. Because of this, inventory costs fell by 12%, production efficiency improved by 10%, and the average stock turnover time got 20% shorter.

Overall, these examples confirm that the reinforcement learning policies proposed by Sultana et al. can be used in real contexts, giving measurable results in lower costs, better product availability, and stronger supply chain resilience.

4.4 AI-based AGV Applications in Inventory Management: The JD.com Case Study⁶⁶

4.4.1 Introduction to JD.com and the competitive environment

JD.com, also known as Jingdong, is today one of the main players in e-commerce in China and at the global level. It was founded in 1998 as an electronics store in Beijing, and in 2004 the company launched its online platform. This move anticipated the digitalization of consumption that, in the following years, would expand very quickly. Time over time, JD.com consolidated its own position and became the second online retailer in China in terms of revenue: now it has millions of active users and a very diversified product portfolio. What makes JD.com different from its competitors, especially Alibaba, is its business model: this model is based on the direct control of the supply chain. While a giant as Alibaba mainly acts as a marketplace where sellers and buyers are connected without managing logistics, instead JD.com has invested heavily in its own infrastructure. The company has its own distribution centres, warehouses and delivery fleets and it takes responsibility for the entire fulfilment process. Even if this strategy required a higher investments at the beginning of its journey, it allow now JD.com to keep a stronger control over operations and to offer a more efficient logistics service.

Currently, the JD.com network includes more than 1,300 warehouses across China, with a total area of more than 30 million square meters. Thanks to this extensive coverage, the company is able to reach not only the biggest cities but also more peripheral and rural regions, offering fast deliveries even in areas that are usually difficult to serve. More than 90% of orders made on JD.com are delivered the same day or within 24 hours, which has helped the company strengthen its reputation among Chinese consumers. These results are possible thanks to a constant focus on technological innovation. The Chinese e-commerce sector is extremely competitive, and events such as Singles' Day (November 11) or the 618 Shopping Festival (June 18) generate exceptional peaks of demand. During these moments, order volumes can increase up to ten times compared to a normal day, creating a strong pressure on warehouse capacity. To keep service standards high even in these situations, JD.com has adopted solutions based on artificial intelligence and robotics, especially the introduction of Automated Guided Vehicles (AGVs) for inventory management.

The use of AGVs responds to very practical needs: reducing internal handling times, improving the accuracy of picking operations, and optimizing the use of space. Automation also reduces the dependence on manual labor, which is a critical factor in contexts where speed is essential. The AGVs are connected to advanced digital management systems, which monitor stock levels in real time and can also predict future demand, allowing a more efficient organization of products inside warehouses. This strategy is part of a broader vision that JD.com describes as "smart logistics." The idea is not only to improve internal efficiency, but also to turn logistics into a real competitive advantage, making the company different from its rivals and building barriers to entry that are not easy to copy. The investment in automation and artificial intelligence should therefore be considered as a strategic choice: it is not only an answer to operational problems, but also a way to support future growth and consolidate the company's leadership in the market.

In conclusion, JD.com can be seen as an important case of how artificial intelligence can be applied to large-scale inventory management. Its combination of physical infrastructure, robotics, and advanced digital systems shows how innovation in logistics can change traditional supply chain models by bringing benefits in terms of operational efficiency but also customer satisfaction.

4.4.2 Technological architecture: AGVs and AI algorithms for inventory management

The technological architecture introduced by JD.com for managing smart warehouses represents one of the most advanced implementations of logistics automation. It does not only rely on the use of Automated Guided Vehicles (AGVs), but also integrates artificial intelligence and operations research, with the aim of transforming internal handling and inventory management into adaptive and predictive processes.

At the center of the model there are the AGVs, mobile robots that transport entire racks of products from storage areas to workstations. This design makes it possible to move from the traditional picker-to-parts paradigm to the parts-to-picker model, where the inventory moves towards the operator. Such a configuration drastically reduces travel times and increases staff productivity, since workers are no longer required to move around inside the warehouse. The real innovation, however, is not only in the physical automation guaranteed by AGVs, but especially in the algorithms that coordinate them. JD.com has formalized the warehouse management problem as a tripartite matching, where three entities must be matched at the same time: the racks to be moved, the AGVs in charge of transport, and the workstations where picking takes place. This problem, which is combinatorial in nature, increases in complexity with the size of the warehouse: hundreds of robots, thousands of racks and dozens of stations create a decision space that is almost impossible to explore with exhaustive methods. To address this challenge, JD.com has developed an intelligent dispatching system that applies hybrid approaches. The problem is divided into two subproblems: the rack–workstation assignment and the AGV–rack assignment. The first is modelled as a maximum flow problem and solved through algorithms derived from the Hungarian method; the second is addressed with linear programming techniques enhanced by cutting planes. This decomposition makes it possible to respect a crucial constraint: each decision cycle must be completed in less than five seconds, otherwise the operational flow would slow down.

The artificial intelligence component is evident in different aspects. First of all, the algorithms are not static: they apply intelligent heuristics and metaheuristics to explore the solution space efficiently, identifying near-optimal configurations quickly. Second, the system uses machine learning for parameter tuning: it adapts decision rules to observed conditions and historical data. Third, the Warehouse Management System (WMS) integrates predictive demand models that anticipate which SKUs will be requested, placing the most critical racks in more favourable positions. These predictive components mean that the system does not simply react, but prepares in advance for expected demand flows.

The complexity of the architecture is further increased by practical constraints. Many racks are double-sided, but for technical reasons only one side can be accessed during each picking operation. Moreover, SKUs can be distributed across several racks, which creates non-trivial decisions: for example, whether it is better to pick from two nearby racks containing partial quantities, or from a single but more distant rack with the full amount. On top of this, the system must manage three different problems: route conflicts between AGVs, the availability of parking areas near the workstations, and the scheduling of battery recharging cycles. The

AI algorithms therefore need to solve allocation, routing, and scheduling problems simultaneously, within a dynamic and constrained environment.

The system architecture is organized into four layers:

1. Integrated management layer, which interacts with external systems (ERP, OMS).
2. Warehouse management layer, which includes the modules for order, inventory, and location management.
3. Intelligent dispatching layer, the core of the architecture, where AI algorithms and predictive models make operational decisions.
4. Device control layer, which translates decisions into executable commands for the AGVs.

This multi-level structure guarantees modularity and scalability: the company can update or replace the decision-making modules without modifying the entire system. In particular, the intelligent dispatching layer represents the “brain” of the warehouse, where real-time data on orders and vehicle status converge. At this level, AI makes it possible to combine the reactive component (immediate dispatching) with the predictive one (anticipation of demand and potential bottlenecks).

The performance of the system has been validated through simulations and operational experiments. Table 4.7 reports the results of a stress test conducted in the Gu'an warehouse, where an automated facility based on AGVs was compared with a conventional facility. The data show that the automated warehouse, even operating with fewer workers (20 compared to 72) and in a much smaller area (21,528 sq. ft. compared to 161,459 sq. ft.), is able to guarantee clearly higher levels of productivity. In particular, each automated workstation processed on average 149 orders per hour, compared to 52 in the conventional system, with a ratio of 2.9:1. The gap is even more significant in the number of items handled per hour and per workstation (435 compared to 75, ratio 5.8:1). Worker productivity also results much higher: 65 items per hour compared to 19, with a ratio of 3.4:1. Finally, space utilization efficiency shows an improvement of more than seven times (1.21 items per sq. ft. compared to 0.17). These results confirm that the integration between AGVs and intelligent algorithms allows not only to reduce the average order fulfilment time and increase throughput, but also to achieve benefits in terms of accuracy and operational sustainability. The use of a smaller area and fewer workers, while still reaching superior performance, demonstrates how AI-driven automation contributes to raising overall productivity while at the same time reducing costs and inventory errors.

Panel A								
	Orders	Items	SKU types	Workstations	Workers	AGVs	Area (sq. ft)	Time (hr)
Automated	8,918	26,086	198	3	20	26	21,528	20
Conventional	18,705	26,950	6,415	18	72	NA	161,459	20
Panel B								
	Orders/hour/workstation		Items/hour/workstation		Items/hour/worker		Items/sq. ft	
Automated	149		435		65		1.21	
Conventional	52		75		19		0.17	
Ratio	2.9:1		5.8:1		3.4:1		7.1:1	

Table 4.7 - Relevant Statistics Gathered During the Gu'an Stress Test

In conclusion, JD.com's technology shows that the efficiency of smart warehouses does not only depend on physical automation, but mainly on the ability to integrate artificial

intelligence and operations research. The use of AGVs is the visible part of a system that finds its strength in predictive models, learning techniques, and optimization algorithms. It is exactly this synergy between hardware and software that enables the Chinese company to transform inventory management into a sustainable competitive advantage.

4.4.3 Implementation challenges and managerial solutions

The introduction of large-scale systems with AGVs and artificial intelligence algorithms, like in the case of JD.com, has not been without problems. Turning traditional warehouses into fully automated centres created technical, organizational and also managerial challenges, which the company tried to solve with different approaches.

One first challenge is about computational complexity. The tripartite matching problem between AGVs, racks and workstations is very hard to solve in short times. In JD's warehouses, a single decision cycle can involve hundreds of vehicles and thousands of racks. The main issue is to produce good quality solutions in just a few seconds, so that the operations are not slowed down. To handle this, JD.com designed a multi-layer system that splits the big problem into smaller subproblems and applies approximate but efficient algorithms.

Another difficulty is scalability. What works in pilot projects has to work also when the order volume becomes huge, for example during Singles' Day or the 618 Shopping Festival. In those moments, the number of orders can multiply in a few hours, creating very high pressure on the logistics network. To avoid bottlenecks, JD.com used a modular design: the different layers (integrated management, warehouse management, intelligent dispatching and device control) can be scaled separately, depending on where the load is higher.

The physical resources also created problems. AGVs move in shared spaces, so they need to avoid collisions and congestion. Racks are sometimes double-sided, and SKUs are distributed on more racks, which makes the decision more complicated: for example, it can be better to choose two close racks with partial quantities, or one rack further away but with the complete order. JD.com managed this by adding special constraints in the dispatching algorithms and by using monitoring systems that check in real time the position and status of each vehicle.

From an organizational point of view, staff involvement was another important issue. The move to automated warehouses required new skills and specific training. JD.com had to find a balance: robots were introduced, but workers had to stay motivated and involved, without feeling they were completely replaced. Human operators still play a complementary role, since they supervise the system, manage exceptions and make sure that the service runs smoothly.

Data management was also critical. The system depends on a constant flow of information coming from orders, the WMS and the AGV sensors. The quality of this data is essential for the AI algorithms to work correctly. For this reason, JD.com set up redundant collection and validation processes, together with analytics platforms that can detect anomalies quickly. Better data quality also improves inventory accuracy, reducing errors in picking or risks of stock-out.

Finally, long-term strategic choices played a key role. JD.com never saw the adoption of AGVs as a single project, but as part of a wider vision of "smart logistics", which also includes drones for deliveries, autonomous vehicles for urban distribution, and demand forecasting algorithms. This systemic view helped justify the big initial investments, and prepared the ground for a fully AI-driven logistics ecosystem.

In conclusion, the challenges faced by JD.com can be divided into three main levels: technical (algorithmic complexity and operational constraints), organizational (training and staff integration) and managerial (scalability, fast deployment, long-term vision). The company's solutions show that adopting AGVs and AI for inventory management is not only about technology, but also about combining innovation, data governance and change management in a holistic way.

4.4.4 Future prospects

The experience of JD.com with AGVs and artificial intelligence gives some important ideas for the future of automated logistics. Using AI-driven solutions is not really an end point, but more the beginning of a longer process. The goal is to reach an even higher level of integration between digital technologies, robotics and daily warehouse management.

One area of development is scalability. The current algorithms already showed that they can solve large-scale problems in times that are acceptable for daily work. But when the logistics network grows bigger and the number of AGVs increases, the system will need more advanced models. Here techniques like reinforcement learning or distributed optimization could help, because they allow the robots to learn movement strategies by themselves. In this way, the system can reduce calculation time and stay more resilient in changing situations.

Another important direction is last-mile delivery. There is already interest in using autonomous vehicles and drones to cover the final part of the supply chain. The good results with AGVs inside the warehouse can be a base to expand AI also outside, with the idea of lowering delivery costs and reaching more areas, including rural regions that are usually more difficult to serve.

Sustainability is also becoming central. Automated systems not only improve productivity, but they can also reduce energy consumption by using space and routes in a smarter way. In the future, optimization algorithms that focus on sustainability could make logistics a sector with less environmental impact, which is in line with global targets of emission reduction.

Finally, data will play an even bigger role. AGVs, sensors and warehouse systems produce a very large amount of information every day. If this data is analysed with big data analytics and predictive models, it can give useful insights for strategic choices. The combination of AI and data-driven decisions will be key to building supply chains that can adapt and improve constantly.

In short, the case of JD.com shows that the future of logistics will probably mean more automation and more distributed intelligence. AGVs are just the first step, and the direction is towards a fully AI-driven ecosystem.

4.5 AI Applications in Last-Mile Delivery: The Canada Post Case Study

4.5.1 Introduction and Problem Context

Last-mile logistics is one of the most critical and expensive parts of the whole supply chain. In some cases, it makes up more than 50% of the total distribution cost. The efficiency of this step is key for customer satisfaction because it is the moment when the product reaches the final customer. With growing competition, driven by the rise of e-commerce and fast delivery

options like next-day or same-day shipping, companies now need to improve both the reliability and predictability of their deliveries.

In paragraph 3.4, the Vehicle Routing Problem (VRP) was presented as the theoretical way to optimize delivery routes. Finding the best route is only part of the problem: indeed knowing the shortest path does not guarantee an accurate prediction of the delivery time. In practice, delivery times are affected by factors such as traffic, driver behaviour, unexpected events, and weather conditions, that are very difficult to include during the planning phase. For this reason, solving the VRP alone is not sufficient to ensure reliable delivery time estimates. It is necessary to combine route planning with a prediction system that learns from historical data and adapts to the conditions of real life scenario.

In this difficult context, artificial intelligence techniques provide a very promising approach. In particular, machine learning makes it possible to use large amounts of historical data to identify recurring patterns and also estimate the expected delivery time: this is all done by taking into account route characteristics, time of day, and environmental conditions. The use of end-to-end models, which learn directly from observations without the need for manual feature design, reduces modelling complexity and leads to more robust predictions.

The case study written by Arthur Cruz de Araujo and Ali Etemad in 2021⁶⁷ presented in this section focuses on the use of deep neural networks for last-mile delivery time prediction: these techniques are based on a large dataset provided by Canada Post, that is the main postal operator in Canada. The dataset contains millions of delivery records collected in the Greater Toronto Area (GTA) over six months: they also are attached to spatial information (origin and destination coordinates), temporal data (departure and delivery timestamps), and weather data (temperature, precipitation, and snow on the ground). The chosen formulation follows the Origin-Destination Travel Time Estimation (OD-TTE) approach, where the goal is to predict travel duration using only the origin and destination coordinates, without access to the actual routes taken.

This approach offers two main advantages: the first one is that it allows the development of a generalizable model that can work well even without having access to high-resolution tracking data. It also reduces issues related to the uncertainty of planned routes, which may vary depending on driver decisions or unexpected events along the way. The ultimate goal is to improve the accuracy of delivery time estimates and to provide predictive tools that support both customer service improvement and internal resource planning.

4.5.2 System Architecture and Data Preparation

To design a last-mile delivery time prediction system it must be used a pipeline capable of reliably collecting, transmitting, and processing large volumes of data. The approach proposed in the Canada Post study is based on a complex cloud architecture organized according to the Internet of Things (IoT) paradigm: thanks to that the authors are allowed to integrate data sources from the field and then to perform an execution of advanced processing through artificial intelligence models.

The reference architecture follows a three-layer model [Figure 4.4]: Perception Layer, Network Layer, and Application Layer.

The Perception Layer includes data collection devices such as handheld scanners and smartphones used by drivers: they record parcel scans at departure from the depot and at delivery. Each event is associated with a timestamp and GPS coordinates. In addition to these operational data, contextual information such as temperature, precipitation, and snow depth is added: these datas come from weather stations located within the target area.

The Network Layer is responsible for transmitting data from field devices to the central system. This communication takes place through commercial mobile networks (4G/5G) that have to ensure near real-time transfer of delivery information. The use of cloud infrastructure makes it possible to scale processing capacity according to data volume and to maintain continuous synchronization across the different system nodes.

Finally, the Application Layer handles data processing and predictive inference. This is the part where the machine learning engine is located and it use the collected data to estimate the expected delivery time. The model is periodically retrained to include new historical data, gradually improving its ability to generalize. The prediction results can be provided both to operational staff, to support resource planning and driver allocation, and to end customers through tracking applications capable of communicating more accurate arrival times.

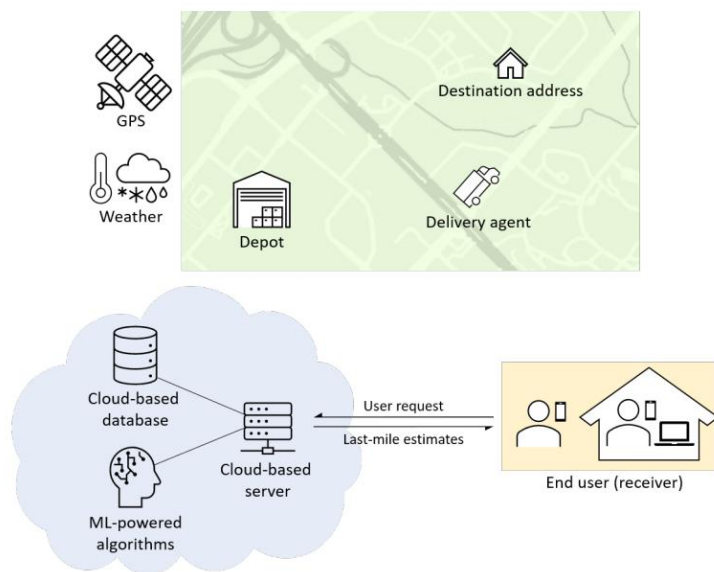


Figure 4.4 – IoT Architecture

Regarding data preparation, the dataset provided by Canada Post includes more than 3.25 million delivery records that were registered in the period January–June 2017 in the Greater Toronto Area. Each record contains the origin (depot) and destination (postal code) coordinates, the departure timestamp, the delivery timestamp, and the weather variables for that day. The geographic coordinates were quantized to a coarser resolution to reduce GPS noise and group nearby locations. In addition, the Haversine distance between origin and destination was calculated and then normalized with respect to the geographic area of the GTA. Several temporal features were extracted from the timestamps, for exaple the hour of the day, the day of the week and the week number. These features were included to capture the variability related to the different times of the day and the seasonal patterns. All numerical variables were normalized in the range $[0,1]$: this was made to ensure stable convergence during the training of the deep learning model. Together, these transformations produce a 12-dimensional input vector for each record, which forms the basis for the modeling phase. This preparation step is essential to ensure the quality of the predictive process. Data consistency, temporal alignment between variables, and correct normalization are key factors that allow the model to learn meaningful patterns without introducing bias or systematic errors.

4.5.3 Deep Learning Models

To address the problem of delivery time estimation, were teste several deep neural network architectures: the goal given tehм was of identifying the model that offers the best trade-off

between predictive accuracy and computational complexity. An end-to-end formulation was adopted: the model takes as input the 12-dimensional vector described in 4.3.2 and outputs the predicted delivery time, without requiring manual feature engineering or prior knowledge of the actual route taken. Three main categories of convolutional architectures were explored: VGG-based networks, ResNet-based networks, and versions of both enhanced with Squeeze-and-Excitation (SE) blocks.

VGG networks are characterized by a sequence of blocks made of two convolutional layers with a ReLU activation that is followed by a max-pooling layer. The goal was to assess whether increasing the depth of the model would improve its ability to capture complex patterns in the data. Some variants with increasing depth, from 3 to 10, were implemented [Table 4.8] The results showed an initial improvement as depth increased, but then was followed by a performance drop beyond seven blocks: it could be due to overfitting or vanishing gradient issues. The following Table summarizes the metrics, highlighting that the VGG-6 variant represents the best compromise, with a Mean Absolute Error (MAE) of 0.8867 hours and a MAPE of 39.36%.

	MSE	RMSE	MAE	EW _{90%}	MAPE	MARE
VGG-3	1.8273	1.3518	0.9789	2.0117	44.04	30.71
VGG-4	1.7542	1.3245	0.9417	1.9527	41.83	29.55
VGG-5	1.7028	1.3049	0.8961	1.8749	39.45	28.12
VGG-6	1.6979	1.3030	0.8867	1.9160	39.36	27.82
VGG-7	1.6743	1.2939	0.9179	1.8864	41.46	28.80
VGG-8	1.6955	1.3021	0.9003	1.8623	40.35	28.25
VGG-9	1.7411	1.3195	0.9363	1.8931	42.17	29.38
VGG-10	1.7494	1.3227	0.9408	1.9258	41.94	29.52

Table 4.8 – Effect of depth in VGG architectures

To overcome the limitations of the deeper VGG architectures, it was tested a second category of models based on Residual Networks (ResNet). These networks introduce skip connections that is a technique that helps to facilitate gradient propagation during training and allow deeper models to be trained without performance degradation. Variants with an increasing number of residual blocks, from three to ten, were implemented. As reported in the following Table 4.8, increasing the depth produced an almost monotonic improvement in the metrics, with ResNet-8 achieving the best results in terms of Mean Absolute Error (MAE = 0.8404 h) and 90% Error Window (EW90% = 1.768 h), making it the preferred choice for system deployment.

	MSE	RMSE	MAE	EW _{90%}	MAPE	MARE
VGG-6	1.6979	1.3030	0.8867	1.9160	39.36	27.82
SE-VGG-6	1.7048	1.3057	0.9003	1.8814	39.86	28.25
ResNet-8	1.5492	1.2447	0.8404	1.7680	36.55	26.37
SE-ResNet-8	1.5516	1.2456	0.8512	1.8292	37.34	26.71

Table 4.9 – Effect of squeeze and excitation augmentation

Finally was evaluated the impact of Squeeze-and-Excitation (SE) blocks. These modules are designed to dynamically recalibrate the importance of feature channels. The hypothesis was that selective attention could further improve data representation and, consequently, predictive performance. However, as reported in the previous Table, the inclusion of SE blocks did not lead to significant improvements and, in some cases, slightly worsened performance.

Model training was carried out using the Adam optimizer, with the loss function defined as Mean Squared Error (MSE). An early stopping strategy was applied if no improvement was observed on the validation set for 25 consecutive epochs, and the initial learning rate was halved every 40 epochs. The models were implemented in Keras with a TensorFlow backend and trained on an Nvidia RTX 2080 Ti GPU, as indicated in the paper.

In summary, after testing different neural network models, the ResNet ones (with skip connections) gave the best results for predicting delivery times. ResNet-8 was the most balanced model: it is deep enough to capture complex patterns but not too big to become slow or hard to train. Other models, like deeper VGG networks or those with Squeeze-and-Excitation blocks, did not bring clear improvements. Overall, ResNet-8 seems the best choice for a delivery time prediction system because it gives good accuracy while keeping the computational cost reasonable. This makes it practical to use in a real logistics environment.

4.5.4 Results and Analysis

The experimental analysis carried out on the Canada Post dataset made it possible to compare different deep learning architectures with a set of traditional baselines, in order to evaluate their effectiveness in last-mile delivery time prediction. The evaluation metrics included Mean Squared Error (MSE), Root Mean Squared Error (RMSE), Mean Absolute Error (MAE), Mean Absolute Percentage Error (MAPE), Mean Absolute Relative Error (MARE), and the 90% Error Window (EW90%). This last metric is very important from an operational point of view, as it indicates the time interval in which 90% of the predicted deliveries fall.

The results reported in Table 4.9 clearly show the superiority of deep learning models compared to traditional benchmarks. For example, the ResNet-8 model achieved a MAE of 0.84 hours: that is a 0.13h reduction compared to values above 0.97 hours for the best traditional method (SB-TTE). This means that the average prediction error for each delivery is reduced by about 10 minutes. On a scale of millions of shipments, this means that there could be a significant reduction of complaints and customer service calls. The EW90% value was also reduced by about 36 minutes if it is compared to the baseline: this means that 90% of deliveries are now predicted with an error of less than two hours. For the final customer, this results in a narrower and more reliable delivery window, increasing the chance of being present when the delivery arrives. From an operational point of view, a smaller uncertainty interval allows the companies to planning better shift and resources: at the same time the can reduce idle times and failed deliveries. In terms of percentage error, the lowest MAPE was of ResNet-8 that recorded a value of 26.37%. By translating this data into real numbers, for a delivery that lasts an average of three hours, the average error is less than 50 minutes. This level of accuracy is difficult to achieve with purely statistical models, which cannot capture complex patterns such as seasonal variations or differences between depots.

	MSE	RMSE	MAE	EW _{90%}	MARE	MAPE
SB-TTE [13]	2.0918	1.4463	0.9728	2.0698	39.34	30.51
ST-NN-TimeDNN [15]	2.3629	1.5372	1.1891	2.2197	56.87	37.31
DNN [17]	2.2649	1.5050	1.1377	2.2032	50.34	35.69
Random Forest	2.4432	1.5631	1.2081	2.2573	57.76	37.90
MLP-1	2.4506	1.5655	1.2147	2.2749	58.07	38.11
MLP-2	2.2940	1.5146	1.1587	2.2017	53.88	36.36
VGG-6	1.6979	1.3030	0.8867	1.9160	39.36	27.82
SE-VGG-6	1.7048	1.3057	0.9003	1.8814	39.86	28.25
ResNet-8	1.5492	1.2447	0.8404	1.7680	36.55	26.37
SE-ResNet-8	1.5516	1.2456	0.8512	1.8292	37.34	26.71

Table 4.10 – Results comparison

In addition to the aggregate metrics, a detailed analysis of the error distribution was carried out to identify possible recurring patterns. Figure 4.5 shows a geographic map of the Greater Toronto Area with the error distribution for each depot. It can be observed that peripheral areas tend to have lower errors, likely due to more regular traffic conditions and the more frequent use of highway segments. In contrast, depots located in downtown Toronto show higher error values, mainly because of heavier urban congestion and variability caused by traffic lights and high-density delivery zones.

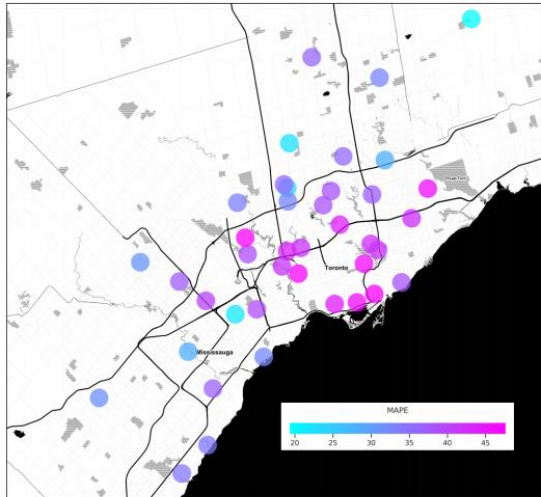


Figure 4.5 - Geographical distribution of MAPE across the 40 busiest depots

The analysis by origin–destination distance shows that predictions are more accurate for distances greater than 7 km [Figure 4.6]. This may seem counterintuitive but can be explained by the fact that longer trips are usually completed on high-speed roads and are therefore more predictable, whereas shorter routes are more affected by local factors such as neighbourhood traffic and frequent stops. From a temporal perspective, Figure 4.7 shows that errors are generally lower for deliveries started early in the morning, when traffic is lighter and operating conditions are more stable. The day-of-week analysis [Figure 4.8] also shows a slight reduction in error on Saturdays and a significant improvement on Sundays, consistent with the lower traffic congestion during weekends. Another perspective is provided by the analysis by delivery duration [Figure 4.9] the models are particularly accurate for deliveries shorter than seven hours, while the error tends to increase as the duration grows. This result is partly related to the definition of MSE, which penalizes deviations more heavily when the target values are high.

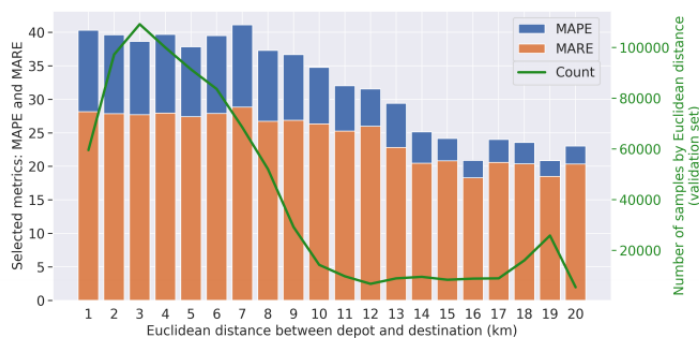


Figure 4.6 - Distribution of MAPE and MARE across the euclidean distance between the depot and the delivery destination

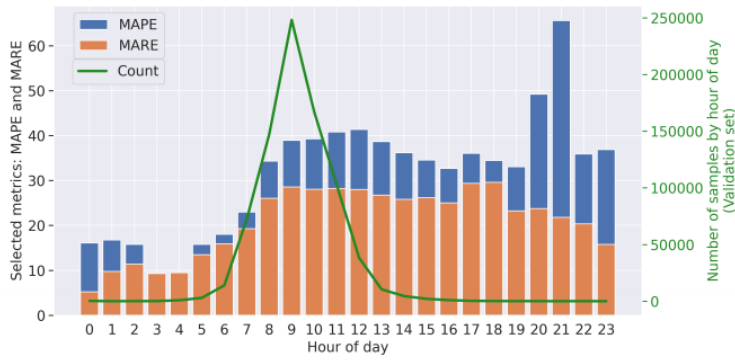


Figure 4.7 - Distribution of MAPE and MARE across the out-for-delivery hour of day

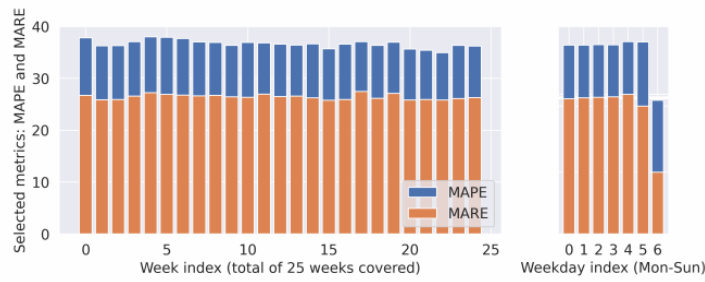


Figure 4.8 - Distribution of the MAPE and MARE across the 25 weeks (left) and 7 days of the week

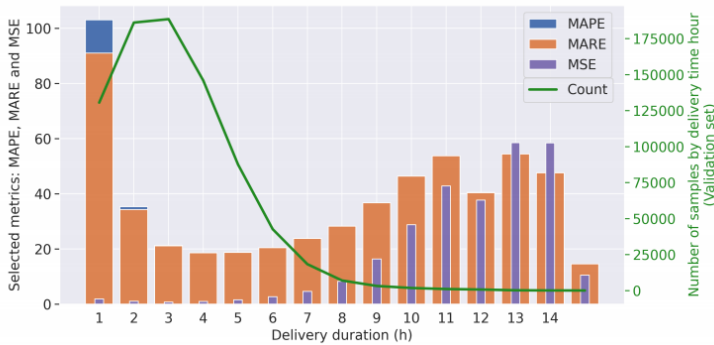


Figure 4.9 - Distribution of the MAPE, MARE, and the MSE across binned hours for the prediction target (delivery duration)

In summary, the experimental results show that the proposed approach not only outperforms traditional models but also provides stable and interpretable predictions across different geographic and temporal scenarios. These findings confirm the validity of the model as a decision-support tool for last-mile logistics management.

4.5.5 Implications and Conclusions

The results show that using deep learning models, especially ResNet architectures, can noticeably improve the accuracy of last-mile delivery time predictions. This has clear effects both for day-to-day operations and for long-term strategy in a logistics company.

On the operational side, having a lower Mean Absolute Error (MAE) and a smaller 90% Error Window (EW90%) makes it easier to plan driver shifts and assign resources more precisely. More accurate forecasts mean narrower delivery windows, which help improve on-time

performance and the chance of succeeding on the first delivery attempt. Less variability in the estimates also cuts waiting time in depots, lowering parking costs and making better use of the fleet.

From a strategic perspective, better prediction accuracy can become a real competitive advantage. The results show that ResNet-8 reduced the error window by about 36 minutes compared to traditional methods, making it possible to give customers more reliable delivery time information. This can improve customer satisfaction and lower the number of complaints and failed deliveries. For companies like Canada Post, which handle millions of shipments, even a 5–10% gain in accuracy can have a major economic impact. The analysis of error patterns also provides useful insights for further local optimization. Indeed, the fact that errors are higher in the urban depots and during the peak traffic hours, suggests the possibility of developing adaptive models: these models should be calibrated for specific geographic areas or certain times of the day. This would make it possible to maximize accuracy where variability is highest, such as in city centres, while keeping simpler models for peripheral areas where traffic is more predictable.

From a technological perspective, the proposed model is light enough to be integrated into a scalable cloud system and used in near real time, allowing predictions to be updated dynamically as new data are collected. Looking ahead, combining the model with live traffic data or with information on the actual route would enable real-time updates for the customer, bringing the estimate even closer to the actual delivery time

In conclusion, using deep neural networks to predict last-mile delivery times seems to be a practical and effective solution. The results from the Canada Post case study show that an end-to-end approach, based only on historical and contextual data, can perform better than traditional models and bring real benefits for both the company and the customer. In the future, these systems could become more dynamic, combining different data sources and giving real-time updates. This would help build supply chains that are more resilient and better focused on customer service.

4.6 Practical Application of AI-based Computer Vision in Quality Control

4.6.1 Case study introduction and goals

Reducing food waste and improving quality control are now key challenges for the agri-food industry. A large part of the production is still lost or thrown away along the supply chain, with serious economic and environmental effects. For this reason, in the last few years several technological systems have been developed to make the selection and classification of fresh products faster, more consistent, and less dependent on the subjective judgment of workers.

The case study in this section looks at an integrated system that uses computer vision and collaborative robotics to automate tomato sorting. The aim of the project is twofold: to detect and classify tomatoes by ripeness and quality, and to move them carefully so they can be separated by category without being damaged. The system was designed to be low cost, but also easy to add to the existing production lines and then to keep human intervention to a minimum (Filho et al., 2024)⁶⁸.

On the hardware side, the system uses a Raspberry Pi 5 with a PiCamera Module 3 to capture images and a UR3e collaborative robotic arm that has a soft gripper based on the Fin Ray Effect. This type of gripper can hold the tomatoes gently, so they are not damaged. The images are processed directly on the Raspberry Pi, which then sends commands to the robot over an

Ethernet connection. This setup is compact and modular, and it keeps the cost low, which makes it suitable even for small and medium farms or processing plants.

For the classification task, two model setups were tested. In the first setup, tomatoes were divided into four classes: unripe, ripe, overripe, and rotten. This gives a very detailed picture of ripeness, but it can be hard to manage, especially when the difference between classes is very small. For this reason, a simpler setup was also tried, grouping tomatoes into just two classes: suitable (unripe and ripe) and unsuitable (overripe and rotten). This simpler approach is so much closer to what is needed in a real sorting plant: there the main goal is to separate tomatoes that can be sold from those that need to be discarded, rather than describing ripeness in detail.

The experiments showed that the model was able to recognize most of the tomatoes correctly. The results were so much better when the images were clear and the fruit was easy to see from the camera. Using the simpler classification made the system more stable, forcing it to make fewer mistakes and more reliable performance for large-scale use. The tests with the robot also showed that the vision system and the robotic arm can work together well: they move and sort tomatoes safely and repeatably, without damaging them. Overall, this case study shows that it is possible to build an integrated and low-cost system for quality control of fresh products like fruits or vegetables. It reduces the variability given by a human inspection, makes the process faster, and helps cut waste along the supply chain.

4.6.2. System Architecture and Experimental Setup

The system architecture was designed to be reliable, easy to integrate, and low cost, so that it can be used even in small and medium production facilities. The solution has three main parts: the image acquisition and analysis subsystem, the robotic manipulation subsystem, and the communication link between them.

Vision subsystem

The recognition part of the system is based on a Raspberry Pi 5 with a PiCamera Module 3. It was chosen because it is compact, it has optimized libraries for image processing, and it uses only a little amount of power. The camera is positioned on a modular frame made of steel and 3D-printed parts, placed vertically above the work area at a fixed distance of about 650 mm. This setup ensures a consistent field of view and an accurate conversion from pixel coordinates to metric coordinates: its very important to have a setup like this because it is necessary for the robot arm to move correctly.

The lighting is provided by two 12V LED projectors, that are positioned to give uniform light and avoid shadows or reflections that could interfere with the segmentation algorithm. Having a good image quality is essential to reduce detection errors and make classification more reliable.

Robotic subsystem

For the manipulation part it was used a UR3e collaborative robot arm. It was chosen because it is flexible, can be connected to external controls, and is safe to use in places where people are working nearby. At the end of the arm there is a soft gripper based on the Fin Ray Effect, which copies the way fish fins bend. This kind of gripper adapts to the shape of whatever it has to grab, so it can hold the tomatoes firmly without damaging them. The gripper was made with 3D printing using flexible material, so it is strong but can still bend when needed.

Communication and coordination

The Raspberry Pi talks to the robot arm over an Ethernet cable. The commands are sent in URScript using Python's socket library. This makes it possible to control the arm in real time,

matching the moment the tomato is detected with the movement of the robot. The algorithm finds the position of the tomato in the image, converts it into real-world coordinates, and sends them to the robot, which moves to that point.

Work area

The whole system is mounted on an aluminium frame, which gives stability and makes it easy to adjust when needed. The work area was defined so that the robot arm has enough space to move freely and the camera can see the entire surface, with as few blind spots as possible. After several tests, a “safe” working area of about 262×250 mm was identified. Inside this area, the robot can do pick-and-place operations repeatedly without risk of collision.

This modular setup makes it possible to adapt the system to different production lines by only changing the camera height or the layout of the work area. This keeps the solution flexible and scalable, so it can be used both in small labs and in larger industrial facilities.

4.6.3 Model Training Methodology

Training the computer vision model was the main step in developing the system, since the ability to correctly identify and classify tomatoes depends on it. A dataset of 7,454 images was created for this purpose, with more than 55,000 manual annotations. The images came from public sources, mainly Kaggle and Roboflow Universe, and were resized to 320×320 pixels to keep the input consistent during training. The dataset was split into three parts: 70% for training, 25% for validation, and 5% for final testing. This split was used to get a reliable estimate of how well the model could work on new, unseen data.

For the architecture, the Single Shot MultiBox Detector (SSD) with a MobileNet v2 backbone was selected. This model works very well with the devices with limited resources, such as the Raspberry Pi 5: it gives a good balance between accuracy and inference speed. The model was trained using the TensorFlow Object Detection API, which includes pre-trained weights that help reduce training time. During training, the model steadily improved, with both classification loss and localization loss decreasing over time. The first metric shows how well the model can tell the classes apart: instead, the second one shows how accurately it can find the tomatoes in the image. The total loss, which combines both, reached low values toward the end, meaning the model had learned a good balance between accuracy and generalization.

After the training phase, the model was converted to TensorFlow Lite, a lighter and optimized version designed to run on embedded hardware. The tests with the converted model showed performance that was good enough for industrial use. The mean average precision (mAP) over the IoU range 0.5–0.95 was 67.54% for the four-class model and 64.66% for the binary model. The best results were seen with large, well-lit images, while accuracy dropped when tomatoes were partly hidden or the image quality was low. Even with these issues, the overall performance was considered good enough for use on an automated sorting line, where the main goal is to separate good tomatoes from the ones that must be discarded rather than classifying every ripeness stage perfectly.

This approach made it possible to build a balanced system, fast enough for real-time processing and accurate enough to meet the needs of a real production environment.

4.6.4 Results and Performance Evaluation

The performance evaluation considered both the accuracy of the computer vision model and its integration with the robotic arm, in order to check whether the system could work properly in a real tomato-sorting scenario.

For the visual recognition phase, the training results showed a steady decrease in classification loss and localization loss, reaching low and stable values after about 10,000 iterations. In the paper, this trend is clearly shown in the training curves [Figures 4.10, 4.11, and 4.12]: they highlight how the model improves its accuracy as training goes on. Regularization loss also decreased over time, helping to avoid overfitting and showing that the model could generalize well to new data.

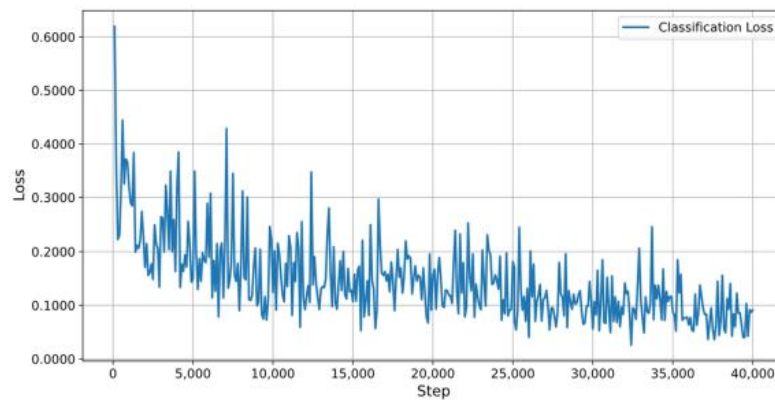


Figure 4.10 -Loss of classification during training—tomato detection and classification model—four classes.

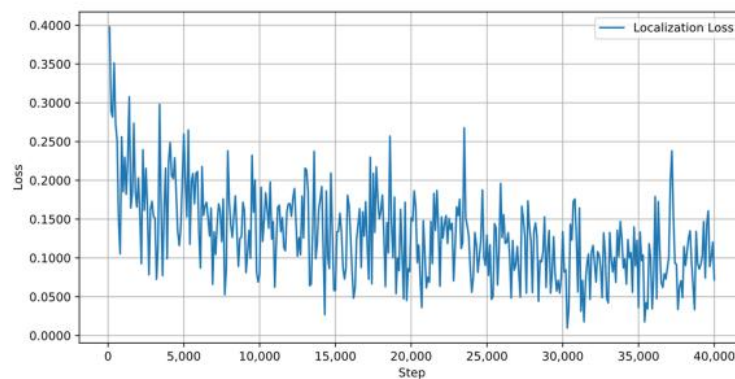


Figure 4.11 - Loss of localization during training—tomato detection and classification model—four classes

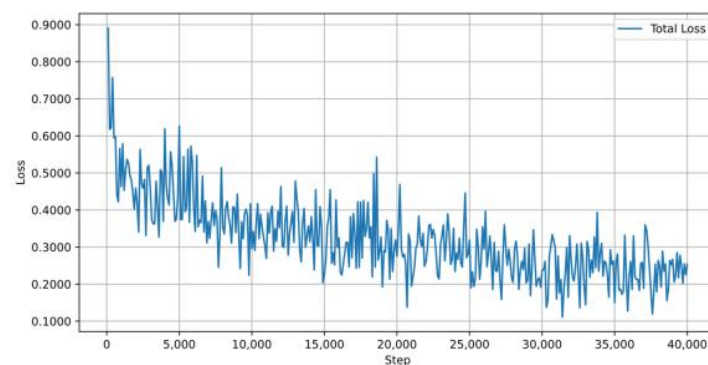


Figure 4.12 - Total loss during training—tomato detection and classification model—four classes

The quantitative evaluation was done using mean Average Precision (mAP), the standard metric for object detection models. The four-class model reached an mAP of 67.54% in the IoU range between 0.5 and 0.95, with very good results on large images, where the average precision went over 77%. Performance dropped on small images, where accuracy was close to

zero, which is expected for models of this type. Table 4.10 in the paper summarizes these results, separating them by image size (small, medium, large) and IoU values.

Metric	IoU	Area	Value
Average Precision (AP)	[0.50:0.95]	All	0.738
	0.50	All	0.911
	0.75	All	0.836
	[0.50:0.95]	Small	0.000
	[0.50:0.95]	Medium	0.416
	[0.50:0.95]	Large	0.774
Average Recall (AR)	[0.50:0.95]	All	0.801
	[0.50:0.95]	Small	0
	[0.50:0.95]	Medium	0.599
	[0.50:0.95]	Large	0.825

Table 4.11 - TensorFlow model results—tomato detection and classification model—four classes

After converting the model to TensorFlow Lite, the performance stayed at a similar level. The mAP was 78.69% at an IoU of 0.5 and gradually dropped as the IoU threshold increased, reaching around 17% at 0.95. Table 4.11 in the paper lists the Average Precision (AP) for each of the four classes: overripe and rotten tomatoes had the best results, both with AP values above 89%, while the unripe and ripe classes showed more variation.

IoU	Tomato Half Ripe AP	Tomato Overripe AP	Tomato Ripe AP	Tomato Rotten AP	mAP
0.5	65.49	90.72	68.82	89.74	78.69
0.55	65.34	90.72	68.55	89.74	78.59
0.6	65.2	89.7	67.35	89.62	77.97
0.65	65.2	89.7	67.03	89.49	77.86
0.7	64.63	87.95	66.67	89.34	77.15
0.75	63.06	87.01	65.93	88.67	76.17
0.8	60.99	87.01	62.72	86.82	74.38
0.85	52.41	82.59	56.11	82.24	68.34
0.9	32.63	65.25	34.17	65.03	49.27
0.95	4.27	25.6	13.54	24.32	16.93

Table 4.12 - TensorFlow Lite model results—tomato detection and classification model—four classes

The simpler two-class model also gave good results, with a slightly lower mAP but better robustness in terms of generalization. The training curves (Figures 4.13, 4.14, 4.15) follow a similar trend to the four-class model, with total loss stabilizing quickly and regularization loss steadily decreasing (Figure 4.16). This simpler setup reduced classification errors in ambiguous cases, making the system more suitable for real industrial use, where the main goal is simply to separate good tomatoes from bad ones.

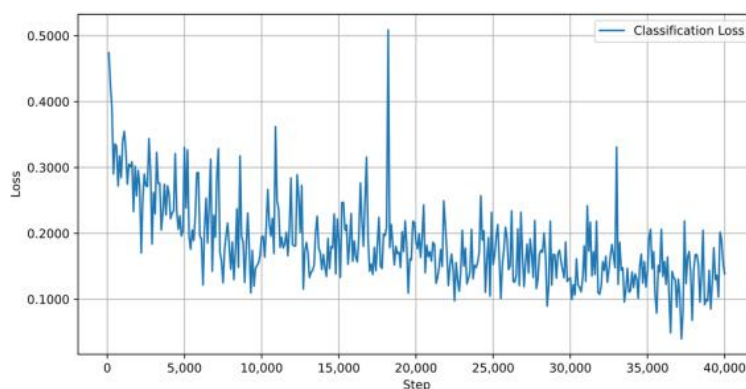


Figure 4.13 - Loss of classification during training—tomato detection and classification model—two classes

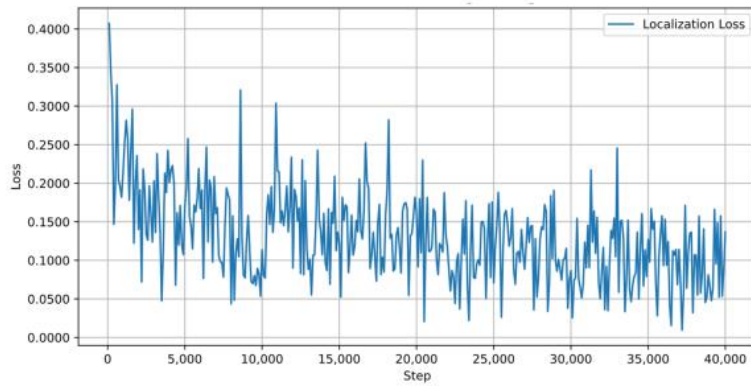


Figure 4.14 - Loss of localization during training—tomato detection and classification model—two classes

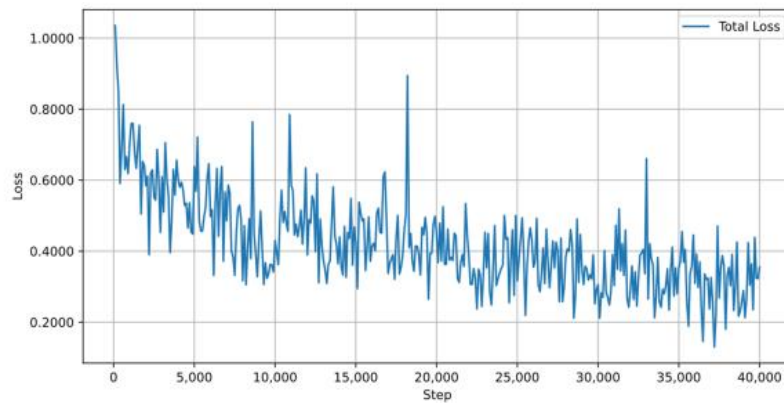


Figure 4.15 - Total loss during training—tomato detection and classification model—two classes

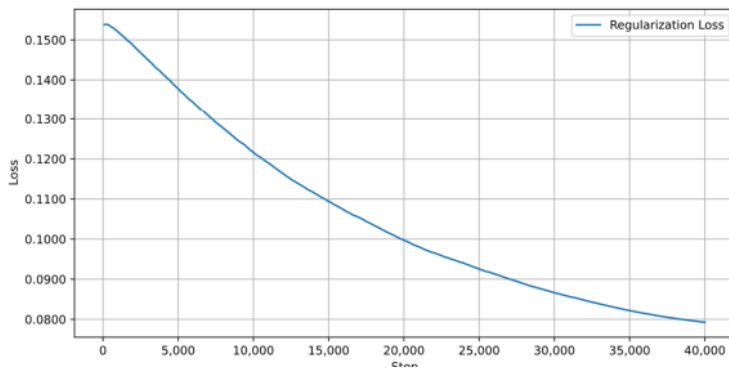
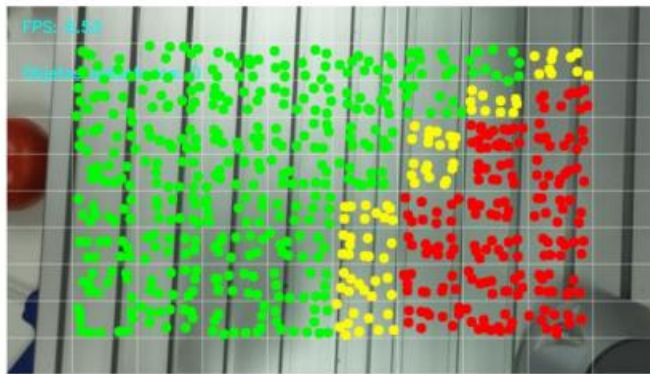


Figure 4.16 - Loss of regularization during training—tomato detection and classification model—two classes.

Finally, 640 manipulation tests were carried out to check the integration between computer vision and the robot. These tests helped define a reliable work area of about 262×250 mm, where the robot arm can perform pick-and-place operations without collisions and with high repeatability [Figure 4.17]. The experiment also showed that the few failures were not caused by the vision algorithm but by mechanical limits of the robot near the edges of the work area.



- Handling task successfully completed
- Manipulation task successfully completed, but kinematics dangerous
- Arm collided into itself or kinematically impossible area

Figure 4.17 - Configuration of the proposed work area - System test results

Overall, the results show that the system can detect and can handle tomatoes with an accuracy level suitable for an automated sorting line. The integration of the vision and robotics proved to be stable and consistent, making possible to scale the system up and adapt it to other types of fruits or vegetables in the future.

4.6.5 Implications for the Supply Chain and Conclusions

This project does not just show that an automatic system based on computer vision can work but also has a bigger meaning for the agri-food supply chain. Being able to classify products in real time, directly on the production line, helps to react faster when quality standards are not met. Lots with a high number of bad products can be found immediately, so the next processing steps can be stopped, and time, energy, and resources are not wasted.

Another important effect is that it makes the process more standard. Manual checks depend so much on the experience of the operators and they can be subjective: instead, an automatic system always applies the same rules. This gives more reliable data and can make it easier to agree on shared quality metrics along the supply chain. Having structured digital data also makes it possible to connect the system with traceability platforms or predictive tools: this feature can help to forecast demand and plan logistics better.

Finally, the system can be adapted for other fruits or vegetables with little extra work. It is enough to collect a new dataset and train the model again with the new classes. This makes the investment more useful in the long term, because the same setup can be reused for different products.

From an operational point of view, using computer vision systems together with collaborative robots does not just reduce the need for manual inspection staff. It also allows workers to be moved to other tasks with higher value, such as process supervision or predictive maintenance. This is important to make sure the technology is accepted in places where human work is still essential.

Looking at the whole supply chain, adopting systems like this helps bring the “Industry 4.0” approach into the agri-food sector, which is usually more traditional. Moving from a reactive quality control to a proactive and digitally connected one creates the base for real quality maintenance strategies, not just defect detection. This change can help reduce waste in a

structural way, predict how much product will be rejected, and in the long run improve both the economic and environmental sustainability of the supply chain.

In conclusion, this case study shows that an integrated system of computer vision and collaborative robotics is not only a promising technology to improve quality control, but also a tool that can drive digital transformation in the supply chain. Its ability to produce reliable data, standardize processes, reduce variability, and give real-time information opens the way to more connected, resilient, and sustainable supply chains.

CHAPTER 5 - From Case Studies to Strategic Insights: The Future of AI in Supply Chain Management

5.1 Integrating Lessons from Real-World Applications

The four case studies in Chapter 4 send a clear message: Artificial Intelligence is no longer just an idea for the future but a real and measurable tool that is already changing supply chains. Each case, Walmart's demand forecasting, reinforcement learning for inventory management, deep learning for last-mile delivery prediction, and computer vision for quality control, shows how AI can solve long-standing problems and also make new ways of working possible.

A first common point is the value of data-driven decisions. The Walmart example showed that moving from basic regression models to more advanced boosting algorithms can make forecasts much more accurate, reducing errors and giving managers better information for planning. In the same way, in the inventory management case, multi-agent reinforcement learning helped coordinate decisions between warehouses and stores: it has found a good balance between having enough stock and avoiding waste that is something very hard to do with fixed rules. Together, these results show that predictive and prescriptive analytics can change supply chains from being reactive to proactive, so that problems can be anticipated instead of just fixed after they happen.

At the same time, these cases show that using AI is not without problems. The Walmart dataset, for example, covered only a short period and a limited area: so, it is not sure if the models would work the same in the long run. Both the demand forecasting and the delivery time prediction systems worked a bit worse in special situations like holidays or peak hours. This shows that models need to be checked and retrained regularly.

Another important lesson is about scalability and adaptability of the system based on AI. The reinforcement learning study showed that policies trained once could still work well in new scenarios, for example, when new products or stores are added, without having to retrain the whole system. This is a key concept for real companies, where the networks and the demand patterns change all the time. In the same way, the Canada Post case showed that deep learning models trained on rich historical data could generalize well to different locations and times, giving good results even in difficult situations like urban traffic.

The computer vision case for tomato sorting also shows that AI is not only about numbers and optimization. By making quality control more consistent and turning it into structured data, these systems can improve traceability and open the way for more digitalization along the supply chain.

Taken together, the four cases give a clear picture: AI gives the best results when it is part of the whole process, from forecasting demand, to managing inventory, to planning deliveries, and making sure quality standards are met.

In the end, the results suggest that AI works best when it is not seen as a separate tool but as part of a bigger plan that combines technology, process changes, and human knowledge. For managers, the message is simple: to get real value from AI, it is important to link it to business goals, invest in good data, and prepare people to work together with intelligent systems.

5.2 - Strategic and Managerial Implications

The case studies are not only about technology; they also show some lessons that are important for managers who want to adopt AI in supply chains. A first clear point is that AI should be

introduced step by step. Starting with simple projects, like demand forecasting, helps companies gain confidence in data-driven decisions and allows employees to get used to new tools. After this first stage, it is possible to move to more complex systems, such as reinforcement learning, which can give suggestions or even automate reordering decisions.

Another important aspect is the role of people and company culture. A common concern among employees is that AI will replace jobs. The cases show that this is not exactly true because AI usually changes jobs instead of eliminating them. For example at Walmart, better forecasts helped reduce stock problems, but anyways workers were still needed to check results, deal with exceptions, and manage suppliers. Instead in the tomato sorting case, computer vision did not remove manual work completely, but it changed tasks of the workers: they are moved towards supervision and maintenance. For this reason, managers need to support workers with training and clear communication, so that AI is seen as a tool to help, not as a threat. Change management, in this sense, is as important as the technology.

Another important aspect concerns fairness, transparency, and data privacy. Companies need to act responsibly in the way they use data, especially in Europe where regulation is strict because of the General Data Protection Regulation (GDPR). Customers and suppliers are becoming more attentive to how their information is collected, stored, and used, and if there is little transparency this can reduce trust and damage long-term relationships. Beyond compliance, there is also a very important reputational factor: organizations that show a clear and ethical data governance are usually considered more reliable partners in the supply chain. On the contrary, weak protection or misuse of data can lead not only to legal penalties but also to financial losses, operational risks, and reputational damage. For these reasons, fairness and privacy should be seen not only as legal requirements but also as strategic elements that influence the competitiveness and resilience of supply chains in the digital era.

More generally, adopting AI should be seen as a strategic change, not only as a technical project. It requires investments in infrastructure, skills, and data governance, but also changes in processes and leadership models. The biggest risk is not that the technology does not work, but that the organization is not ready to the transition. Companies that see AI as a simple plug-and-play solution may be disappointed. Instead, firms that integrate AI into a broader vision, connect it with business goals, and involve people at all levels are more likely to gain long-term benefits.

5.3 - Challenges and Barriers to AI Adoption for SMEs

One of the main points that came out from the study of Artificial Intelligence in supply chains is that not all companies can adopt these technologies in the same way. Big multinational firms can invest a lot of money in infrastructure, experts, and experiments, while small and medium-sized enterprises (SMEs) face a very different situation. For this reason, the enthusiasm about AI should always be considered together with the real difficulties that many firms, especially SMEs, experience when they try to move from theory to practice.

The first barrier is data. Many SMEs still work with fragmented systems, often based on spreadsheets in Excel or old ERP software with little or no integration. This makes data not only limited but also inconsistent, and therefore difficult to use for training reliable models. While large firms have years of detailed information, SMEs often lack both the volume and the variety of data. A second problem is infrastructure. Running AI systems requires environments where data can be processed, models can be trained, and results monitored. SMEs often use outdated servers or local systems that were never designed for this hard transition. Cloud solutions can help SMEs, but they also bring new costs: also they require a digital maturity that is not always present in a small or medium company. Skills are another issue. Experts such as

data scientists or AI engineers are very expensive and hard to find. For a medium-sized firm, hiring such profiles is usually not sustainable. Even when external consultants are used, there is the risk of becoming dependent on them without building internal knowledge. In many cases, the cultural gap is the most important problem, compared to the technical one: adopting AI means changing processes, decision-making, and sometimes the company culture itself. Costs also matter. AI is not only about the first investment but also about maintenance, updates, and retraining of models. For the SMEs, who have usually small margins as they don't reach economy of scale, this can be a real limitation. Finally, there are legal and ethical aspects. Data privacy and compliance with regulations are becoming more complex. Large corporations usually have teams that deal with these problems, while SMEs often do not.

To summarize all, adopting AI in SMEs is possible, but not as easy as for a multinational company. It requires being aware of the barriers and taking a gradual and realistic approach. Otherwise, there is the risk of presenting AI as a universal solution, while in reality its accessibility is still uneven.

5.4 - Roadmap for AI Implementation in Supply Chains

One important lesson from both the literature and the case studies is that the use of AI in supply chains cannot just happen without preparation. Many times companies start with too much expectation, and sometimes they become sceptical after projects that do not work well. For this reason, a clear plan or roadmap is needed, with steps that go one after the other, trying to balance what is possible with what is ambitious.

A roadmap for the adoption of Artificial Intelligence in supply chain management can usually be divided into three main stages. These stages correspond to different levels of maturity of companies. The division is not fixed, because every company must adapt it to its own situation, but it helps to understand how to move from first experiments to full integration.

The first stage is the foundation. In this phase the most important aspect is data. If data is not reliable and well organized, even the most advanced model cannot give good results. For this reason, companies must improve the way they collect and store information and make sure that data can be used by different departments. For small and medium enterprises this step is already very difficult, because many still use old systems or separate tools. However, this step is necessary. Creating a culture where decisions are made on data and not only on intuition is the starting point for every use of AI. The second stage is the experimental phase. In this phase companies can start with simple applications, such as sales forecasting or demand analysis. The goal is not only to obtain first benefits but also to allow employees to see how the models work and to start trusting them. Pilot projects are useful to test the technology but also to help people learn. It is very important to involve all employees, show them the results and explain how the models work, so that resistance is reduced and acceptance is higher. The third stage is scaling and integration. In this phase AI is no longer an isolated project but it is part of the supply chain strategy. Applications are extended, and AI is integrated in decision-making at different levels. More advanced systems, such as reinforcement learning or digital twins, can be used to optimize inventory, simulate different scenarios, or even automate some activities. The objective is to move from simple support to partial automation, but always with human control to guarantee strategic alignment and ethical responsibility.

Of course, this roadmap is not the same for all, indeed every company has to adapt it to its own resources, culture, and objectives. The important point is to move step by step, always connecting technology with business goals, and keeping a balance between innovation and sustainability. It is also necessary to manage expectations by workers, managers but also

shareholders or owners: AI can be very powerful, but it is not magic. Without a clear plan, it can easily be seen as only a temporary trend.

In conclusion, a roadmap is not just about technology. It is also about organizational change. It means creating the right conditions of data, skills, culture, and strategy so that AI can really help supply chains to become stronger and more competitive.

5.5 - Future Outlook and Research Directions

The future of Artificial Intelligence in the supply chain management will probably move from single, isolated projects to a more integrated use of these technologies. What today looks like many separate applications will in the next years become more connected, with forecasting, optimization, and decision-making working together along the whole supply chain.

A key step in this direction will be the use of digital twins⁶⁹ together with the Internet of Things. Traditional systems usually give only a partial or static picture, while digital twins can create a virtual copy of the supply chain. With the real-time data coming from IoT sensors, these models can be updated constantly and used for many applications: for example, to test strategies, simulate possible problems, and see bottlenecks before they really happen. Beyond simulation, digital twins also make it possible to evaluate “what-if” scenarios in a controlled environment, reducing the risks of experimenting directly on operations. They allow managers to anticipate disruptions and compare alternative courses of action, integrating predictive and prescriptive analytics into daily decision-making. This makes the system more dynamic compared to the traditional planning methods and helps supply chains to become more flexible but also adaptive and resilient in the face of uncertainty.

Another important development is the move from predictive models to prescriptive or even autonomous systems. Today AI is often used only to forecast demand or find inefficiencies. In the future, it will also suggest concrete actions, and in some cases take them automatically if certain conditions are met. Reinforcement learning, for instance, is already tested for reordering policies and could in the next years be used for semi-automatic procurement. This change does not mean that humans will disappear from the process. Instead, their role will change: managers will spend less time on routine operations and more time on supervising, validating results, and making sure that AI decisions are aligned with the company’s strategy. Human–AI collaboration will therefore be a key point, not only for technology but also for organization and governance.

Another important issue that firms should be aware of in the future, is the robustness of the models that they will use. While such models generally provide satisfactory results under standard operating conditions, their performance tends to deteriorate in the presence of atypical scenarios such as sudden demand peaks, unexpected market shocks, or disruptions linked to geopolitical crises and extreme weather events. These contexts are more and more frequent and they highlight the need to develop approaches that do not only optimize average-case accuracy but are able to ensure stability and reliability in highly volatile environments. Several research directions appear promising in this respect. Adaptive algorithms represent a first option, as they are capable of recalibrating parameters continuously on the basis of new data streams, thereby maintaining prediction accuracy even when operating conditions change rapidly. Transfer learning constitutes another line of investigation, allowing models trained on large datasets in well-documented contexts to be reused and adapted in situations where historical data are scarce, with clear advantages in terms of generalization and cost reduction. A further possibility is the design of hybrid systems in which machine learning models are combined with rule-based mechanisms, so as to exploit the capacity of statistical methods to detect complex patterns together with the reliability of predefined procedures in critical situations. The

integration of these approaches could lead to predictive systems that remain effective not only in stable conditions but also in the presence of external shocks, thereby increasing the resilience of supply chains and supporting continuity of operations.

Another area that will change the future of AI is the sustainability behind this technology. Supply chains must, starting from today, balance efficiency with environmental and social goals. Going in this direction, AI will not be used only to cut costs or improve lead times but also to reduce emissions, optimize energy consumption, and support more ethical sourcing practices. But at the same time, the environmental impact that AI itself will have on the planet can't be ignored. For example, the energy used to train OpenAI's GPT-4⁷⁰ has been estimated as enough to power 50 American homes for 100 years. This shows that very large models can become a sustainability problem on their own. To face this issue, researchers and tech companies are looking at alternative energy sources for data centres. One option that has been proposed is the use of small modular nuclear reactors (SMRs) to provide stable and low-carbon power for cloud infrastructures that run advanced AI systems. Even if this solution is technically possible, its real adoption will depend on costs, regulation, and public acceptance in the next years.

The adoption of AI will not proceed in the same way for all companies. Small and medium-sized enterprises (SMEs) generally face greater constraints, such as limited financial resources, smaller and less structured datasets, legacy IT systems, and difficulties in attracting or retaining specialized professionals. If these obstacles are not addressed, there is a concrete risk that AI will widen the gap between large corporations and smaller firms. This is not only an economic issue but also a social one, since the resilience of supply chains depends on the participation of a broad base of suppliers and not just on the largest actors. Nonetheless, SMEs are not excluded from this transition. Several strategies can support their gradual involvement. Cloud-based platforms and "software as a service" solutions reduce the need for significant infrastructure investments, while collaborations with technology providers, research centres, or universities can give access to skills that would otherwise be too costly to develop internally. Just as important is adopting a progressive approach: starting with limited and low-risk applications, such as demand forecasting or inventory monitoring, makes it possible to demonstrate concrete benefits and build confidence in the technology before moving on to more advanced systems. By doing this, even smaller firms can incorporate AI into their processes and maintain competitiveness, thereby contributing to the overall resilience of supply chains.

In conclusion, the future of AI in supply chains will not depend only on technology. It will also depend on organization, rules, and society. AI must be robust, transparent, sustainable, and fair. Its development will not be a finished project but a continuous process, where data, models, and people work together. The companies that will benefit most will be those that see AI not as a single tool but as part of a bigger strategy, with technology, governance, and culture together.

ACKNOWLEDGMENT

BIBLIOGRAPHY AND WEBSITE

- ¹ McCarthy, J., Minsky, M. L., Rochester, N., & Shannon, C. E. (2006). A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence, August 31, 1955. AI Magazine.
- ² Russell, S., & Norvig, P. (1995). Artificial intelligence: A modern approach. Prentice Hall.
- ³ Trend Micro, *Cos'è l'Intelligenza Artificiale?*
- ⁴ Bailey, K. (2016, October 27). Reframing the “AI Effect”. Medium.
- ⁵ Turing, A. M. (1950). Computing machinery and intelligence. Mind.
- ⁶ McCarthy, J., Minsky, M. L., Rochester, N., & Shannon, C. E. (2006). A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence, August 31, 1955. AI Magazine.
- ⁷ Newell, A., & Simon, H. A. (1956). The logic theory machine: A complex information processing system. RAND Corporation.
- ⁸ MYCIN, Encyclopaedia Britannica.
- ⁹ Rivista.AI, L'inverno dell'IA (dal 1974 al 1980), 7 luglio 2024.
- ¹⁰ Mitchell, T. (1997). Machine Learning. McGraw-Hill Education.
- ¹¹ Cortes, C., & Vapnik, V. (1995). Support-vector networks. Machine Learning. Kluwer Academic Publisher.
- ¹² Goodfellow, I., Bengio, Y., & Courville, A. (2016). Deep Learning. MIT Press.
- ¹³ Stanford CS229: Machine Learning Course, Lecture 1 - Andrew Ng (Autumn 2018).
- ¹⁴ Goodfellow, I., Bengio, Y., & Courville, A. (2016). Deep Learning. MIT Press.
- ¹⁵ The Economist. (2024, September 19). The breakthrough AI needs. The Economist. Retrieved June 22, 2025, from The Economist website.
- ¹⁶ MIT Introduction to Deep Learning | 6.S19.
- ¹⁷ Jurafsky, D., & Martin, J. H. (2023). *Speech and Language Processing* (3rd ed.).
- ¹⁸ MIT OpenCourseWare, Lecture 1: Introduction to Machine Vision.
- ¹⁹ Siciliano, B., & Khatib, O. (Eds.). (2016). *Springer Handbook of Robotics* (2nd ed.). Springer.
- ²⁰ AIX Digest. (2025, circa marzo). How AI-driven smart factories are transforming the future of manufacturing. LinkedIn.
- ²¹ Council of Supply Chain Management Professionals. (2011). Definition of Supply Chain Management.
- ²² Ballou, R. H. (2004). Business logistics/supply chain management (5th ed.). Upper Saddle River, NJ: Pearson Education.
- ²³ Chopra, S., & Meindl, P. (2016). Supply Chain Management: Strategy, Planning, and Operation (6th ed.). Pearson Education.
- ²⁴ Chopra, S., & Meindl, P. (2016). Supply Chain Management: Strategy, Planning, and Operation (6th ed.). Pearson Education.
- ²⁵ SAP. . What is a warehouse management system (WMS)?
<https://www.sap.com/products/scm/extended-warehouse-management/what-is-a-wms.html>.
- ²⁶ SAP. . What is a transportation management system (TMS)? Retrieved July 7, 2025.
<https://www.sap.com/products/scm/transportation-logistics/what-is-a-tms.html>.
- ²⁷ ediBasics. (n.d.). EDI vs API: Definition & advantages. <https://www.edibasics.com/edi-vs-api/>
- ²⁸ Hashemi-Pour, C., & Sutner, S. (2024, 24 ottobre). What is a Business Intelligence Dashboard (BI Dashboard)? TechTarget.
<https://www.techtarget.com/searchbusinessanalytics/definition/business-intelligence-dashboard>.
- ²⁹ HighRadius. (n.d.). What is Electronic Invoice Presentment & Payment (EIPP)?
<https://www.highradius.com/resources/Blog/eipp/>
- ³⁰ Modern Treasury. (n.d.). What is a treasury management system (TMS)?
<https://www.moderntreasury.com/learn/what-is-a-treasury-management-system/>
- ³¹ Lee, H. L., Padmanabhan, V., & Whang, S. (1997). "The Bullwhip Effect in Supply Chains." *Sloan Management Review*.
- ³² Priyesh, S. (2024, June 4). How TikTok, Supply Chains and Viral Demand Impact Forecasting. Inside Supply Management. <https://ominthenews.com/tiktok-shops-impact-on-supply-chains/>
- ³³ Seuring, S., & Müller, M. (2008). From a literature review to a conceptual framework for sustainable supply chain management. *Journal of Cleaner Production*.
- ³⁴ Il Sole 24 Ore. (n.d.). Ok finale Parlamento europeo – stop vendita auto inquinanti 2035.
<https://www.ilsole24ore.com/art/ok-finale-parlamento-europeo-e-stop-vendita-auto-inquinanti-2035-AESigHnC>.
- ³⁵ Wijffelaars, M., & van Harn, E.-J. (2022, 12 aprile). Ukraine war poses a threat to EU industry. RaboResearch, Rabobank.

<https://www.rabobank.com/knowledge/d011280916-ukraine-war-poses-a-threat-to-eu-industry>

³⁶ Davidson, H. (2025, April 16). China trade war: Beijing targets US arms firms by curbing rare earths supply. *The Guardian*.

³⁷ Mui, C. (2025, February 02). What Trump's trade war means for 'friendshoring'. *Politico*.

³⁸ Maersk. (2025, February 28). Navigating the shifting geopolitical supply chain landscape. Maersk Insights.

³⁹ Center for Strategic and International Studies. (2023). Crossroads of commerce: How the Taiwan Strait propels the global economy. CSIS.

⁴⁰ Cohen, M. C., & Tang, C. S. (2024, February 5). The role of AI in developing resilient supply chains. *Georgetown Journal of International Affairs*. <https://gjia.georgetown.edu/2024/02/05/the-role-of-ai-in-developing-resilient-supply-chains/>.

⁴¹ Gartner. (2023). *IT Glossary*. Gartner. <https://www.gartner.com/en/information-technology/glossary>

⁴² Fattah, J., Ezzine, L., Aman, Z., El Moussami, H., & Lachhab, A. (2018). Forecasting of demand using ARIMA model. *International Journal of Engineering Business Management*. <https://journals.sagepub.com/doi/full/10.1177/1847979018808673>.

⁴³ Croston, J. D. (1972). Forecasting and stock control for intermittent demands. *Operational Research Quarterly*.

⁴⁴ Vandeput, N. (2019, January 25). Croston forecast model for intermittent demand. Medium – Data Science.

⁴⁵ Pastor, E. (2024–2025). *Business intelligence per big data*. Lecture slides. Master's Degree Program in Management Engineering – Politecnico di Torino.

⁴⁶ Cortes, C., & Vapnik, V. (1995). Support-vector networks. *Machine Learning*.

⁴⁷ Smola, A. J., & Schölkopf, B. (2004). A tutorial on support vector regression. *Statistics and Computing*.

⁴⁸ Pastor, E. (2024–2025). *Business intelligence per big data*. Lecture slides. Master's Degree Program in Management Engineering – Politecnico di Torino.

⁴⁹ Goodfellow, I., Bengio, Y., & Courville, A. (2016). *Deep learning*. MIT Press.

⁵⁰ DataMasters. (2024, July 26). Reti neurali artificiali: cosa sono e come funzionano. DataMasters Blog.

⁵¹ Rumelhart, D. E., Hinton, G. E., & Williams, R. J. (1986). Learning representations by back-propagating errors. *Nature*.

⁵² Carboneau, R., Laframboise, K., & Vahidov, R. (2008). Application of machine learning techniques for supply chain demand forecasting. *European Journal of Operational Research*.

⁵³ Hopp, W. J., & Spearman, M. L. (2011). *Factory physics: Foundations of manufacturing management* (3rd ed.). McGraw-Hill Education.

⁵⁴ Talbi, E.-G. (2009). *Metaheuristics: From Design to Implementation*.

⁵⁵ Gen, M., & Cheng, R. (2000). *Genetic Algorithms and Engineering Optimization*. Wiley-Interscience.

⁵⁶ Kallenberg, L. C. M. (2010). *Markov Decision Processes*. Lecture Notes, University of Leiden.

⁵⁷ Sutton, R. S., & Barto, A. G. (2018). *Reinforcement Learning: An Introduction* (2nd ed.). MIT Press.

⁵⁸ Toth, P., & Vigo, D. (2014). *Vehicle Routing: Problems, Methods, and Applications* (2nd ed.). SIAM.

⁵⁹ Perboli, G., Tadei, R., & Vigo, D. (2011). The Two-Echelon Capacitated Vehicle Routing Problem: Models and Math-Based Heuristics. *Transportation Science*.

⁶⁰ A MIP solver (Mixed Integer Programming solver) is an algorithmic solver used to solve mathematical programming problems involving both continuous and integer variables. These solvers combine techniques such as branch-and-bound, branch-and-cut, and relaxation methods to efficiently explore the solution space and find the optimal or approximate solution within a given time limit.

⁶¹ Yang, J. (2024). A hybrid method combining reinforcement learning and heuristics in solving two-echelon vehicle routing problem with backhauls. In Z. Chen, X. Zhao, & Y. Xu (Eds.), *Advances in artificial intelligence: Algorithms and applications* (pp. xxx–xxx). Springer. https://doi.org/10.1007/978-981-97-5495-3_18

⁶² Ettalibi, A., Elouadi, A., & Mansour, A. (2024). AI and computer vision-based real-time quality control: A review of industrial applications. *Procedia Computer Science*.

⁶³ Neba, C., Gerard, S. F. B., Nsuh, G., Amouda, P., Neba, A., Webnda, F., Ikpe, V., Orelaja, A., & Sylla, N. A. (2024). Advancing retail predictions: Integrating diverse machine learning models for accurate Walmart sales forecasting. *Asian Journal of Probability and Statistics*. <https://doi.org/10.9734/ajpas/2024/v26i7626>

⁶⁴ Sultana, N., Gosavi, A., & Das, T. K. (2020). Reinforcement Learning for Multi-Product Multi-Node Inventory Management in Supply Chains. arXiv preprint arXiv:2006.04037. Disponibile su: <https://arxiv.org/pdf/2006.04037>

⁶⁵ SuperAGI (2025). Real-World Success Stories: How Top Companies Are Optimizing Inventory with AI in 2025. Disponibile su: <https://www.superagi.com/real-world-success-stories-how-top-companies-are-optimizing-inventory-with-ai-in-2025>

- ⁶⁶ Qin, H., Wei, Z., Fang, L., & Zhang, D. (2022). JD.com: Operations research algorithms drive intelligent warehouse robots to work. *INFORMS Journal on Applied Analytics*, 52(1), 42–55. <https://doi.org/10.1287/inte.2021.1100>
- ⁶⁷ de Araujo, A. C., & Etemad, A. (2021). *End-to-end prediction of parcel delivery time with deep learning for smart-city applications*. arXiv preprint arXiv:2009.12197. <https://arxiv.org/abs/2009.12197>
- ⁶⁸ Filho, E. V., Lang, L., Aguiar, M. L., Antunes, R., Pereira, N., & Gaspar, P. D. (2024). Computer Vision as a Tool to Support Quality Control and Robotic Handling of Fruit: A Case Study. *Applied Sciences*, 14(21), 9727. <https://doi.org/10.3390/app14219727>
- ⁶⁹ Schuster, R., Mitjavila, L., & Penazzo, C. (2024, July 29). Using digital twins to manage complex supply chains. Boston Consulting Group. <https://www.bcg.com/publications/2024/using-digital-twins-to-manage-complex-supply-chains>
- ⁷⁰ The Economist. (2024, 19 settembre). The breakthrough AI needs. The Economist. <https://www.economist.com/leaders/2024/09/19/the-breakthrough-ai-needs>.