

CHAPTER 1

Demand-Driven Forecasting in the Supply Chain

The world is changing at an increasing pace. Consumers are becoming more demanding, and they expect products and services of high quality, value for their money, and timely availability. Organizations and industries across the globe are under pressure to produce products or provide services at the right time, quantity, price, and location. As global competition has increased, those organizations that fail to be proactive with information and business insights gained risk loss of sales and lower market share. Supply chain optimization—from forecasting and planning to execution point of view—is critical to success for organizations across industries and the world. The focus of this book is on demand-driven forecasting (using data as evidence to forecast demand for sales units) and how cloud computing can assist with computing and Big Data challenges faced by organizations today. From a demand-driven forecasting perspective, the context will be a business focus rather than a statistical point of view. For the purpose of this book, the emphasis will be on forecasting sales units, highlighting possible benefits of improved forecasts, and supply chain optimization.

Advancements in information technology (IT) and decreasing costs (e.g., data storage, computational resources) can provide opportunities for organizations needing to analyze lots of data. It is becoming easier and more cost-effective to capture, store, and gain insights from data. Organizations can then respond better and at a quicker pace, producing those products that are in high demand or providing the best value to the organization. Business insights can help organizations understand the sales demand for their products, the sentiment (e.g., like or dislike products) that customers have about their products, and which locations have the highest consumption. The business intelligence gained can help organizations understand what price sensitivity exists, whether there is effectiveness of events and promotions (e.g., influencing demand), what product attributes make the most consumer impact, and much more. IT can help organizations increase digitalization of their supply chains, and cloud computing can provide a scalable and cost-effective platform for organizations to capture, store, analyze, and consume (view and consequently act upon) large amounts of data.

This chapter aims to provide a brief context of demand-driven forecasting from a business perspective and sets the scene for subsequent chapters that focus on cloud computing and how the cloud as a platform can assist with demand-driven forecasting and related challenges. Personal experiences (drawing upon consultative supply chain projects at SAS) are interspersed throughout the chapters, though they have been anonymized to protect organizations worldwide. Viewpoints from several vendors are included to provide a broad and diverse vision of demand-driven forecasting and supply chain optimization, as well as cloud computing.

Forecasting of sales is generally used to help organizations predict the number of products to produce, ship, store, distribute, and ultimately sell to end consumers. There has been a shift away from a push philosophy (also known as inside-out approach) where organizations are sales driven and push products to end consumers. This philosophy has often resulted in overproduction, overstocks in all locations in the supply chain network, and incorrect understanding of consumer demand. Stores often have had to reduce prices to help lower inventory, and this has had a further impact on the profitability of organizations. Sales can be defined as shipments or sales orders. Demand can include point of sales (POS) data, syndicated scanner data, online or mobile sales, or demand data from a connected device (e.g., vending machine, retail stock shelves). A new demand-pull (also known as an outside-in approach) philosophy has gained momentum where organizations are learning to sense demand (also known as demand-sensing) of end consumers and to shift their supply chains to operate more effectively. Organizations that are changing their sales and operations planning (S&OP) process and moving to a demand-pull philosophy are said to be creating a demand-driven supply network (DDSN). (See Figure 1.)

The Boston Consulting Group (BCG) defines a demand-driven supply chain (DDSC) as a system of coordinated technologies and processes that senses and reacts to real-time demand signals across a network of customers, suppliers, and employees (Budd, Knizek, and Tevelson 2012, 3). For an organization to be genuinely demand-driven, it should aim for an advanced supply chain (i.e., supply chain 2.0) that seamlessly integrates customer expectations into

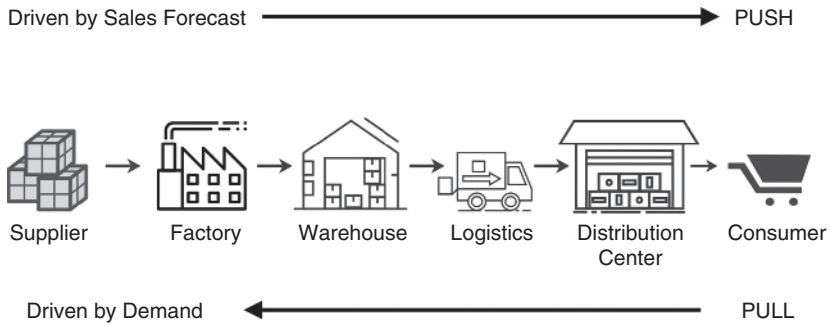


Figure 1 Push and Pull—Sales and Operations Process

its fulfillment model (Joss et al. 2016, 19). Demand-driven supply chain management focuses on the stability of individual value chain activities, as well as the agility to autonomously respond to changing demands immediately without prior thought or preparation (Eagle 2017, 22). Organizations that transition to a demand-driven supply chain are adopting the demand-pull philosophy mentioned earlier. In today's fast-moving world, the supply chain is moving away from an analog and linear model to a digital and multidimensional model—an interconnected neural model (many connected nodes in a mesh, as shown in Figure 2). Information between nodes is of various types, and flows at different times, volumes, and velocities. Organizations must be able to ingest, sense (analyze), and proactively act upon insights promptly to be successful. According to an MHI survey that was published (Batty et al. 2017, 3), 80 percent of respondents believe a digital supply chain will prevail by the year 2022. The amount of adoption of a digital supply chain transformation varies across organizations, industries, and countries.

It has become generally accepted that those organizations that use business intelligence and data-driven insights outperform those organizations that do not. Top-performing organizations realize the value of leveraging data (Curran et al. 2015, 2–21). Using business intelligence (BI) with analytics built upon quality data (relevant and complete data) allows organizations to sense demand, spot trends, and be more proactive. The spectrum of data is also changing with the digitalization of the supply chain. Recent enhancements in technologies and economies of

capture and storage, and the Internet of Things (IoT) to transform their business to a digital supply chain (a well-connected supply chain). Such data and analytics can lead to improved insights and visibility of an entire supply chain network. The end-to-end supply chain visibility of information and material flow enables organizations to make holistic data-driven decisions optimal for their businesses (Muthukrishnan and Sullivan 2012, 2). Organizations wishing to optimize their supply chain management are moving toward an intelligent and integrated supply management model that has high supply network visibility and high integration of systems, processes, and people of the entire supply chain network internal and external to the organization (Muthukrishnan and Sullivan 2012, 2–5).

The holistic and real-time data coupled with advanced analytics can help organizations make optimal decisions, streamline operations, and minimize risk through a comprehensive risk management program (Muthukrishnan and Sullivan 2012, 5). The value of data is maximized when it is acted upon at the right time (Barlow 2015, 22). The benefits of the increased visibility and transparency include improved supplier performance, reduced operational costs, improved sales and operations planning (S&OP) outcomes, and increased supply chain responsiveness (Muthukrishnan and Sullivan 2012, 6). Implementing a supply chain with high visibility and integration provides benefits such as increased sales through faster responses and decision making, reduced inventory across the supply chain, reduced logistic and procurement costs, and improved service levels (Muthukrishnan and Sullivan 2012, 11).

The increasing needs for supply chain visibility are leading to the adoption of supply chain control towers (SCCTs), depicted in Figure 3. An organization could use an SCCT as a central hub to centralize and integrate required technologies, organizations (intranet and extranet supply chain network members), and processes to capture, analyze, and use the information to make holistic and data-driven decisions (Bhosle et al. 2011, 4). Using an SCCT can help with strategic, tactical, and operational-level control of a supply chain. Having a holistic view through an SCCT helps an organization and its supply chain network to become more agile (e.g., ability to change supply chain processes, partners, or facilities). It also helps increase resilience against unexpected events outside of the control of the supply chain network.

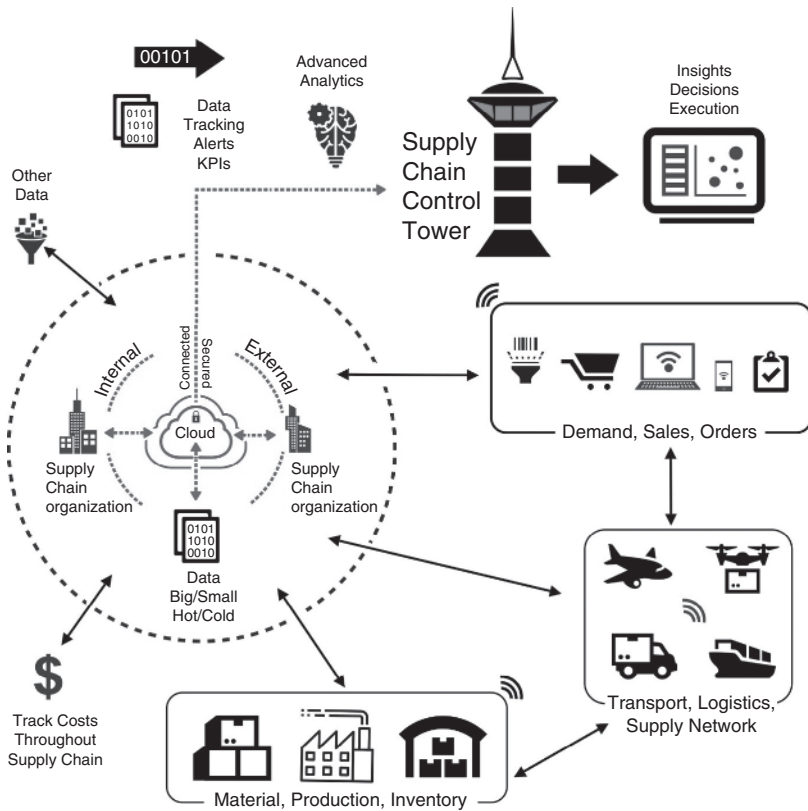


Figure 3 Supply Chain Control Tower

Reliability and supply chain effectiveness can be improved by meeting service levels, cost controls, availability, and quality targets (Bhosle et al. 2011, 4–6).

An SCCT can also help a supply chain network become more responsive to changes in demand, capacity, and other factors that could influence business (Bhosle et al. 2011, 6). There are three phases of maturity for implementing and executing such a supply chain control tower. The first phase typically focuses on operational visibility such as shipment and inventory status. Phase 2 is where the information flowing to the supply chain control tower is used to monitor the progress of shipments through the various network nodes of a supply chain and alert decision makers of any potential issues

or events. In the third and most mature phase, data and artificial intelligence are used to predict the potential problems or bottlenecks (Bhosle et al. 2011, 5–8). The data captured and processed by the SCCT can provide the supply chain visibility and insights necessary to make appropriate decisions and to operate a customer-focused supply chain (Bhosle et al. 2011, 9).

Benefits of a supply chain control tower include lower costs, enriched decision-making capabilities, improved demand forecasts, optimized inventory levels, reduced buffer inventory, reduced cycle times, better scheduling and planning, improved transport and logistics, and higher service levels (Bhosle et al. 2011, 11).

One of the main challenges of the digital supply chain is demand-driven forecasting, and it is generally a top priority of organizations wishing to improve their business. Forecasting and Personalization were ranked as the top two needed analytical capabilities (Microsoft 2015, 14). The forecasting function was rated as either very challenging or somewhat challenging (39 and 36 percent, respectively) in an MHI Annual Industry Report (Batty et al. 2017, 9), and in a 2018 survey more than 50 percent of respondents noted the forecasting function as very challenging (see Figure 4).

There are distinct phases of maturity for forecasting, and such maturity levels vary significantly across organizations, industries, and countries. Unscientific forecasting and planning (e.g., using personal

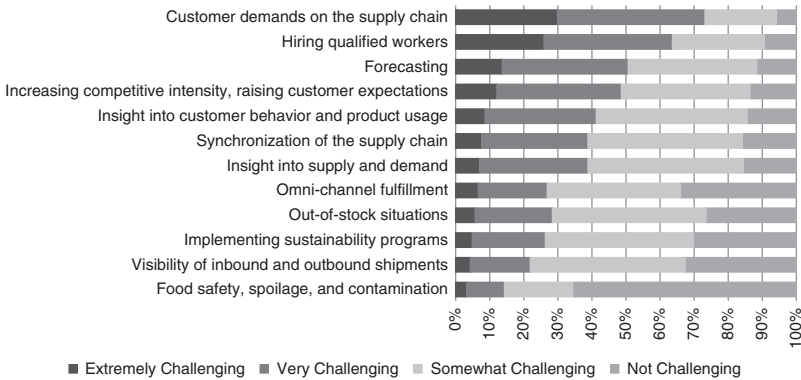


Figure 4 MHI 2018 Survey Results: Company Challenges

Source: MHI Annual Industry Report, 2018, 8.

judgment versus statistical evidence) are still prevalent in many sectors, as shown in a survey by Blue Yonder (2016) in the grocery retail sector. The Blue Yonder report highlights the finding that 48 percent of those surveyed are still using manual processes and gut feeling to make choices, instead of using data-driven actions (Blue Yonder 2016, 25). There are many benefits of making a transition to a demand-driven supply chain. Research by BCG highlights that some companies carry 33 percent less inventory and improve delivery performance by 20 percent (Budd, Knizek, and Tevelson 2012, 3).

A strategy for improved forecasting needs to be holistic and to focus on multiple dimensions to be most effective. The journey toward improvement should include three key pillars:

1. Data
2. Analytics
3. Collaboration—people and processes using a collaborative approach

1. DATA

As mentioned earlier, data is the foundation for analytics, business intelligence, and insights to be gained. The famous “garbage in, garbage out” concept equally applies to today’s challenges. Organizations must be able to capture and analyze data that is relevant to forecasts and supply chain optimizations. Having access to holistic data (e.g., historical demand data, data from other influencing factors) allows organizations to apply advanced analytics to help sense the demand for their products. Insights gained from analytics allows organizations to detect and shape demand—for example, the most demanded products at the right location, at the right time, at the right price, and with the right attributes. Leveraging data and advanced analytics allows organizations to understand correlations and the effect that influencing factors such as price, events, promotions, and the like have on the demand of sales units. As Marcos Borges of the Nestlé organization noted (SAS Institute press release, October 12, 2017), a differentiating benefit of advanced forecasting is the ability to analyze holistic data (multiple data variables) and identify factors influencing demand for

each product throughout a product hierarchy. This process should be automated, and be able to handle large volumes (e.g., many transactions across many dimensions) with depth of data (e.g., a hierarchy of a product dimension).

Quality of data is an essential but often overlooked aspect of analytics. Generally, for a forecast to be meaningful, there should be access to at least two years of historical data at the granularity level of the required forecast (e.g., daily or weekly data for weekly forecasts). This data should be available for all hierarchy levels of the unit or metric of the time series. For example, a consumer packaged goods (CPG) company wishing to predict demand for chocolates would have a product dimension in its data mart for forecasting. This dimension would have a hierarchy with various categories and subcategories. Individual products are called leaf member nodes, and they belong to one hierarchy chain. Those products therefore have a direct and single relationship link rolling upward through the hierarchy. A leaf member can just roll up through one subcategory and category (see Figure 5). Ideally, data should be available for all relevant dimensions. Granular data for the levels of all dimensions should also be available. The combination of product dimension data in this example and time-series data (e.g., sales transactions) that is complete (e.g., sales transaction data across all levels of product hierarchy for at least two years) increases the accuracy of the forecast.

If data is available across all levels of the hierarchy of the dimension, then forecast reconciliation techniques (performed by software solutions) such as top-down, bottom-up, and middle-out forecasting

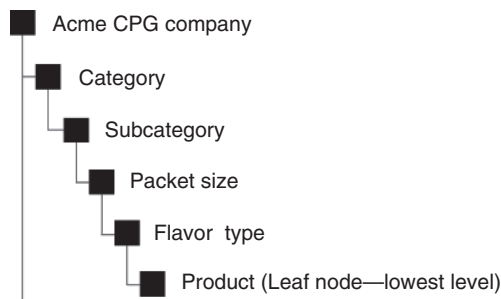


Figure 5 Example: Product Dimension Hierarchy

lead to more accurate results. Ideally, the system would highlight which levels of the hierarchy would provide the most substantial results. These reconciliation techniques aggregate data up or down a hierarchy, allowing forecasts to be consolidated and grouped at various levels. The aforementioned methods can help with demand planning (e.g., using the consolidated forecasted demand at a subcategory or category level). The more data there is available at the granular level (lower levels of the hierarchy of product dimension in this example) the more accurate the aggregation and proportioning can be. Using these methods, a demand planner can then view forecasts at a category level, store level, or regional level, for example.

Typically, other dimensions used in demand forecasting include store location and customers, and these are commonly represented in a star schema data model (see Figure 6). Such a design that separates data can help with the performance of the analytics process used for generating forecasts. The method of striking a balance between all data stored together and separating data is referred to normalization and denormalization of a data model. The data schema design has a profound impact on the analytic capabilities and the performance (speed of completion) of the computations. Therefore, it is equally important to collect the right data (data about metric to be forecasted, as well as data from causal variables), with data of at least two years' time horizon, and to organize the data appropriately (e.g., data marts, logical data schemas). Advancements in data storage and analytic technologies such as data lakes and Big Data can help provide more flexibility and agility for this design process and is elaborated on later in this section.

It is typical for individual analytical functions within supply chain optimization to have separate data marts. For example, data stored for demand forecasts can be stored in one data mart, whereas data related to inventory optimization or collaborative demand planning could each have separate data marts. Such data marts should have easy data integration and allow data flow between functions to enhance the usability of data and increase analytical value. This single or integrated set of data marts for the supply chain analytics is also referred to as a demand signal repository (DSR). (See Figure 7.)

The data model design of these data marts and their storage methods are well suited for advanced analytics (such that demand-driven

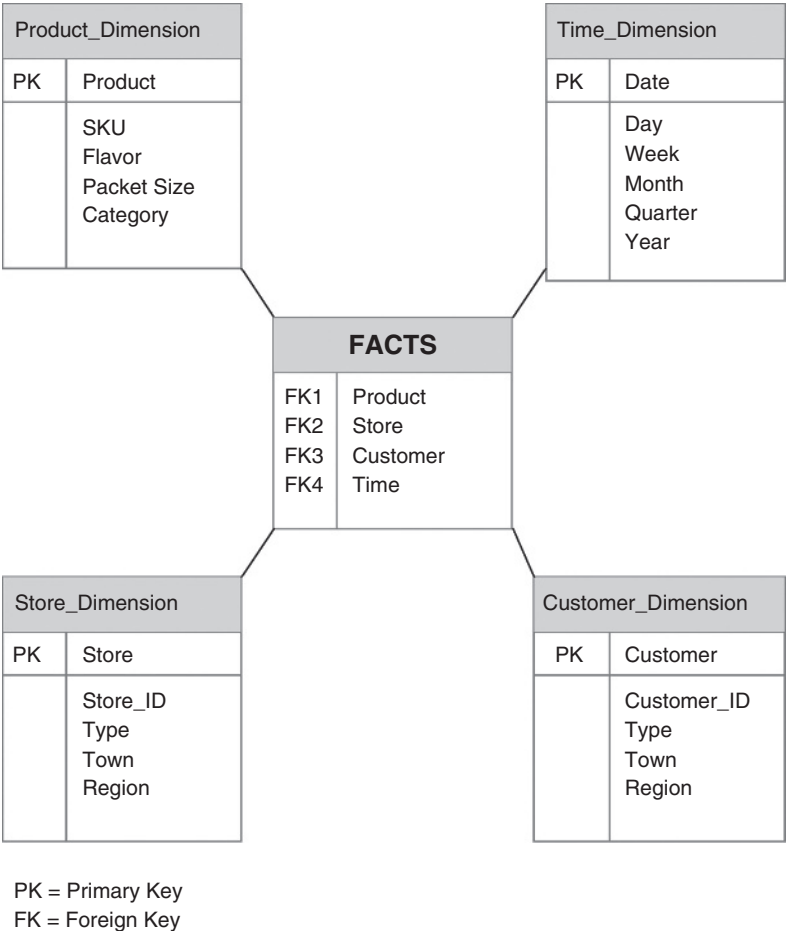


Figure 6 Example: Star Schema - Forecast Dimensions

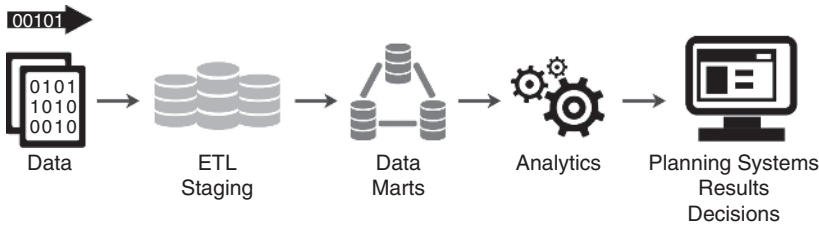


Figure 7 Traditional Data Flow—Supply Chain Analytics

forecasting uses). Business subject matter typically organizes data marts—for example, a data mart for forecasts, a data mart for inventory optimization, a data mart for finance, and so forth.

Slowly moving data (e.g., daily, weekly, or monthly ingress of data) is generally captured via traditional technologies such as databases or data warehouses. Such slowly moving data is also referred to as cold or warm data. Fast-moving data (e.g., per second, minute, hour, day) is captured with the help of connected devices (e.g., IoT), processing technologies such as event stream processing (ESP) that provide near real-time analytics, and advanced data storage (e.g., data lake) technologies and formats that focus on the rapid ingestion of data. This faster-moving data is also referred to as hot data. Real-time analytics with ESP will become a vital component of a connected supply chain in the future. Such data can flow into a supply chain control tower, and can help the organization gain insights from data and act upon it proactively—for example, analyzing data from logistic providers and demand data in real time, and reacting faster to changes in demand or logistics.

Organizations typically use a data mart purpose-built for analytics such as demand forecasting. Separation for a purpose generally increases performance (e.g., separating write operations into online transaction processing [OLTP], and read operations into online analytical processing [OLAP]). Such isolation also allows data management processes to be tailored to the types of data, as well as to the speed of data ingestion. The data storage types can help with analytical loads; for example, OLAP systems are purpose-built for analytics (e.g., searching through data, filtering, grouping, calculations, view multiple dimensions, etc.).

In a digital supply chain, there are many different data sources, which generate different types of data (e.g., point of sales data, retail sales data, online sales data, shipment data, logistical data, etc.). The volume, variety, and velocity of data challenge traditional systems and methods for storing and analyzing data. The data lake is a method to help with these challenges. Organizations can collect data from many sources and ingest these into their data lakes (on premises or in the cloud).

One of the differences of a data lake as opposed to a data warehouse or data mart is that with a data lake organizations initially do not have to worry about any particular data schema (organization of data), nor concern themselves with data transformations at the ingestion stage.

Traditional databases, data warehouses, and data marts follow a schema-on-write approach. Therefore data ingestion processes must extract, transform, and load (ETL) (the overall process is also referred to as data wrangling) the data to fit a predefined data schema. The data schema can include multiple databases or data marts but requires data to match definitions (e.g., data types, data lengths, data stored in appropriate tables). The ETL process is generally defined once or does not change that often, and is most likely scheduled after that. The method of data ingestion, staging data (importing into required formats, storing in a commonly accessible form and location), and ETL can take minutes or hours depending on the complexity of tasks and the volume of data. For example, forecasts at a weekly time granularity level often ingest incremental data at a weekly time interval. Forecasting systems often perform this data import process in batch jobs in nonbusiness hours (e.g., weekend or night time). The forecasting process can also be automated, as can the sharing of data with downstream planning systems. If a demand planning process involves human review and collaboration, then that process is included in a forecast cycle. Depending on the forecast horizon (e.g., time periods into the future) and other factors such as lead times (e.g., supplier, manufacturing) and speed of turnover of the products to be forecasted, the overall forecasting cycle can take hours, days, weeks, or longer.

A data lake follows a schema-on-read approach. In this case the data ingestion process extracts and loads data into a storage pool. Data transformations are performed at a later stage if required (ELT). The data remains in a data lake and can be accessed directly. It can also be transformed and copied to other data storage targets (e.g., databases, data marts), or accessed and leveraged via other means (e.g., data virtualization mechanisms). Such a data ingestion process via a data lake permits fast ingress of data and is typically aimed at fast-moving (hot) data. Analytics on such fast-moving data can occur as quickly as data is ingested. This process is often referred to as event stream processing (ESP) or streaming analytics and focuses on data that is moving in the

second or minute time frames. Using the combination of a data lake and ESP, for example, it is possible to detect values near real time to trigger an event or to traffic light an anomaly.

Tamara Dull, director of emerging technologies at SAS, defines a data lake as follows: “A data lake is a storage repository that holds a vast amount of raw data in its native format, including structured, semi-structured, and unstructured data. The data structure and requirements are not defined until the data is needed” (Dull 2015). A data warehouse stores structured data, has a defined data model that information is molded to (also referred to as a schema-on-write concept), and is mature. This traditional method of ingesting, storing, and using data has a high consistency of data and is used by business professionals for deriving insights.

A data lake, in contrast, can be used to ingest and store structured, semistructured, and raw data. Structured data examples include comma-separated values (CSV) files with defined fields, data types, and order. Semistructured data examples include JavaScript Object Notation (JSON) file (defined fields, ordering, and data types can change). Raw or unstructured data examples include media files (e.g., JPEG, video), emails, documents, and the like. A data lake follows a schema-on-read method, eliminating the need for data wrangling and molding at ingestion time. A data lake is therefore well suited for fast ingestion of data from all types of sources (e.g., streaming, internet, connected devices, etc.). Data lakes are designed to be horizontally scalable, with commodity hardware providing an excellent cost-to-performance ratio for organizations. The maturity of data lake systems is steadily enhancing and, as the use by organizations worldwide and across industries increases, so do the solutions for easy access to data and analytics against such systems.

One of the standard technologies behind a data lake is the Hadoop distributed file system (HDFS) and the Hadoop framework. This technology allows storing any data type on an interconnected grid of computer nodes, leveraging cheaper local storage present in each node. The file system manages the complexity of distributing and managing data files, including redundant copies of each file for high availability (HA), disaster recovery (DR), and performance of computing and analysis (used with methods like MapReduce). The

Hadoop framework leverages cheaper commodity server hardware and scales out horizontally (adding more server nodes as storage or computing requirements dictate). This is a fundamental difference from the traditional framework of vertically scaling servers (increasing computing resources—central processing unit [CPU] and random-access memory [RAM]). The cost of vertically scaling is a lot higher, and, although advancements in computing are continuing, the vertical scale approach has a limit at some point, whereas in theory the horizontal scaling approach has no limit.

Another benefit of this horizontal scaling approach is that data and computing power can stay together on interconnected nodes. A big analytical processing job is broken down into smaller segments (referred to as the mapping phase) and each node in a Hadoop server farm (also called clusters) analyzes segments of that job, based on data that is available to its local storage. The results from each node are then consolidated into one result (this step is the reduce phase).

Once data has landed in a data lake, it can be analyzed, or processed further to be transferred into different formats and a data mart, for example. Simple MapReduce methods enable data to be mined on a large scale and at faster speeds than were previously possible. (See Figure 8.)

EXAMPLE

EXAMPLE

There are many data sets with patient records, and these data sets are distributed across many computer nodes. Typically, there are three copies of each data set, which are hosted on different nodes, assisting with disaster recovery goals. A user would like to report on the number of males per age group across all the data. The user submits a MapReduce job to filter for males per age group from each data set stored across the HDFS, and then consolidate the results. The inner workings of the MapReduce process and the Hadoop framework are out of scope for this book—the aim is to highlight the storage, processing, scalability, and speed (wall clock time) benefits of using a data lake and distributed computing power.

The technologies that enable a data lake, and a data lake itself, can help with the challenges of Big Data. The National Institute of

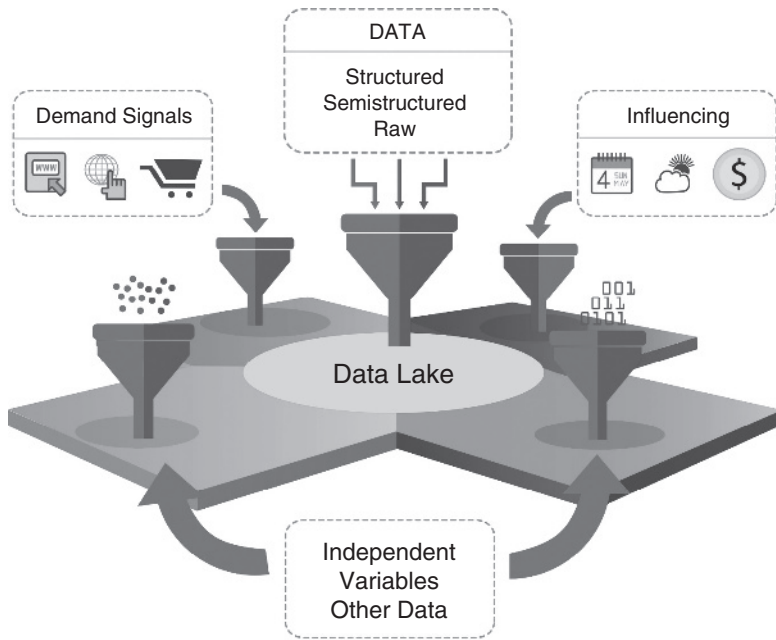


Figure 8 Data Lake - Data for Demand Forecasting

Standards and Technology (NIST) defines Big Data as follows: “Big Data consists of extensive datasets—primarily in the characteristics of volume, variety, velocity, and variability—that require a scalable architecture for efficient storage, manipulation, and analysis” (NIST.SP.1500-1). The Big Data concept can be broken down into two interrelated concepts that need to be addressed if organizations can successfully leverage such technologies. The first concept is the challenge of the data (also known as the 4-Vs). The second concept deals with a change in architecture to enable the 4-Vs of the data.

The 4-Vs of Big Data

1. Volume (i.e., the size of the data to be ingested—could be one or multiple data sets)
2. Variety (i.e., different data types, various data sources, different data domains)

3. Velocity (i.e., the speed of ingestion—could be in seconds, minutes, hours, days, weeks, etc.)
4. Variability (i.e., unexpected change in other characteristics)

Source: NIST.SP.1500-1.

The 4-Vs of Big Data have driven new architectural designs leveraged by IT systems to meet these modern challenges. A modern architecture referred to as the lambda architecture aims to separate functions and layers, enabling a scalable architecture with many components. Such components can perform tasks (e.g., storage, processing, analytics, presenting) on their own, in sequence, or in parallel. The building blocks of such an architecture depend on the software vendor and could be proprietary, open source, or a mixture of both. Extensive details of such architectures are beyond the scope for this book, but at an elevated level, the standard layers of an architecture design following principles of the lambda architecture are depicted in Figure 9. These layers enable ingestion of hot and cold (fast- and slow-moving) data. The processing of data can be sequential or in parallel.

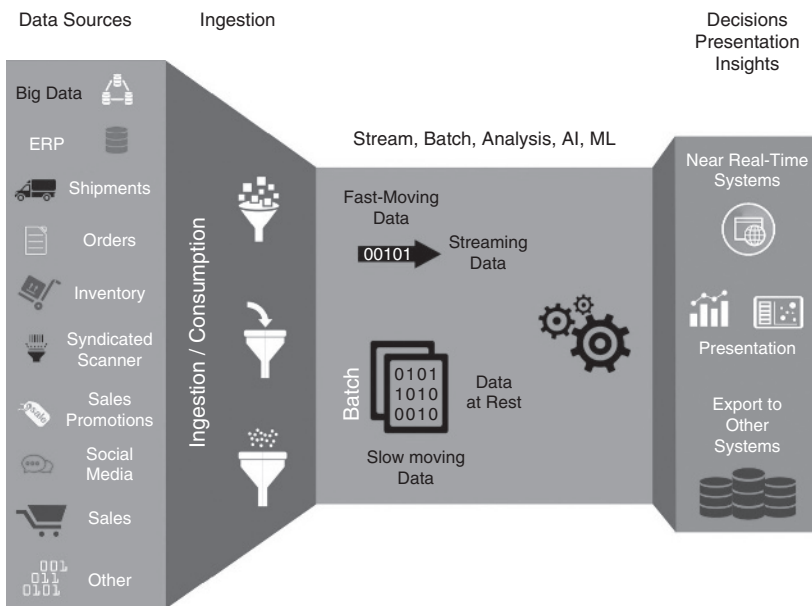


Figure 9 High-level Lambda Architecture Design

Data could be stored in a data lake or sent to different storage and analytics platforms such as databases, data warehouses, and data marts. The analytics layer can also be sequential and in parallel, and handle hot or cold data. The results can be shared with other targets, such as near real-time decision-making processes, different storage platforms, and different systems, or presented as results. This architecture is well suited for a hybrid approach of leveraging hot and cold in-streaming data, as well as already stored data. Analytical processes can combine newly ingested data with previously stored data to provide near real-time results or result sets for further analysis by other systems and processes. This hybrid approach assists with the challenges of the 4-Vs of Big Data. Data ingestion can follow schema-on-write or schema-on-read, leverage different storage systems and data types, and leverage distributed computational resources to provide results aiding data-driven insights promptly. The logical building blocks of a lambda architecture are depicted in Figure 9. Data sources are examples only, based on demand-driven supply chain needs.

Case Study: Leeds Teaching Hospital

To improve its hospital services, it was necessary for Leeds Teaching Hospital to identify trends through vast amounts of data. The primary challenge was the enormous volume of structured and unstructured data. One of the objectives of the health care provider was to detect possible outbreaks of infectious diseases as early as possible using data-driven insights and business intelligence. Previously such analysis relied on cold data that was already stored or archived, and hence out of date. There were enormous insights in text files from a mix of data sources such as unscheduled visits to the accident and emergency (A&E) rooms, retail drug sales, school attendance logs, and so on. Such data could help provide better data-driven insights near real time. The expected volume of data was half a million structured records and about one million unstructured data records. Leveraging data from various sources would provide better insights but would require a lot of computing power. It was not feasible to provision a server farm (lots of server computers) to handle such analysis. Costs, maintenance, and management of the computing environment would be too high a cost of ownership.

(Continued)

(Continued)

The health care provider decided to explore a cloud-based strategy. This would be cost-effective (the hospital would pay only for what it consumed) and would provide scalability and other benefits. Microsoft Azure cloud was chosen as it offered an integrated and seamless stack of solutions and components required (e.g., data ingestion, data lake, processing, business intelligence, presentation, and collaboration), and is one of the leading providers of public clouds in the world. This cloud environment enabled the on-demand processing of six years of data with millions of records. Using a combination of Microsoft's data platform technologies (i.e., SQL Server, HDInsight—a unique Hadoop framework), it was possible to process large volumes of structured and unstructured data. The integration of Microsoft business intelligence (BI) tools enabled a self-service approach to data-driven insights and evidence-based decisions. The digitalization of processes (e.g., data collection, coding, and entry into systems) saved time and reduced stationery and printing costs by a conservative estimate of £20,000 per year. The cloud platform and business model made it possible to spin up a Microsoft Azure HDInsight cluster to process six years' worth of data in just a few hours and shut down the cluster when the analytic job was complete.

Source: Microsoft (September 7, 2014)

In the context of demand-driven forecasting for a supply chain, a hybrid approach could help solve new challenges of the supply chain. Such an approach could combine features of data processing (hot, cold), data storage, analytics, and sending results to downstream systems for decisions near real time or at slower rates. (See Figure 10.)

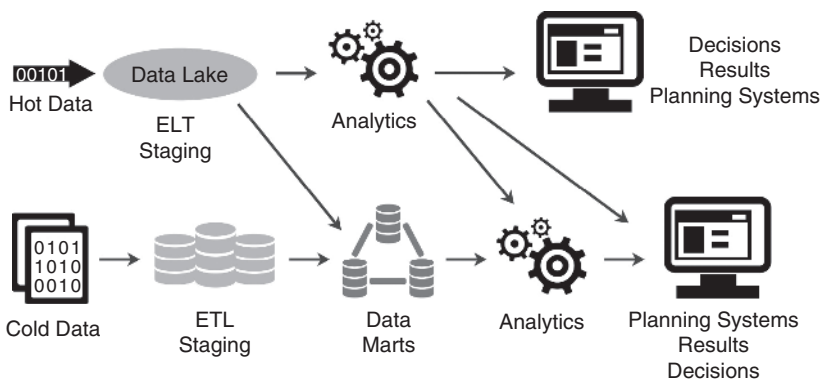


Figure 10 Hybrid Modern Data Flow—Supply Chain Analytics

Leveraging a mixture of data storage, databases, data marts, and analytical technologies is referred to as polyglot persistence. Data virtualization is also a useful technology for abstracting data sources and layers. Data remains at its location, and the data virtualization layer unifies and adds business-friendly access, metadata, and labels.

2. ANALYTICS

While data is an essential foundation toward the result, it is the analytics that provide the highest value for a demand-driven supply chain. Details of statistics, forecast models, and modeling are beyond the scope for this book. The aim is to highlight enhancements and possibilities made possible with cloud computing, and how the combination of disciplines can enhance value further. There are multiple challenges for the analytics of demand forecasting. First, there is the challenge of Big Data (the volume of data that needs to be ingested and analyzed at various speeds). Second, there is the challenge of multiple variables and identifying causal (influencing) variables. Third, there is the challenge of discovering patterns, trends, and links. Such analysis is helpful for detecting changes in consumer behavior and tastes. It is also useful to new product forecasting that can leverage patterns, and information about similar products to assimilate the demand for a new product based on similar attributes and possible tastes. Finally, there is the challenge of automation and leveraging a vast repository of forecasting models and modeling techniques to increase accuracy and value of forecasts. This becomes even more important with multiple dimensions and the depth of those dimensions (the depth of the hierarchy, e.g., tens or hundreds of thousands of products). All these computations must also be time relevant. This could mean near real time, or at least fast enough to fit into a demand forecasting and demand planning cycle and processes.

There are distinct phases of maturity when it comes to analytics, and the envisioned end state an organization wishes to be in will drive the state of advanced analytics leveraged. The four phases are depicted in Figure 11 and are also called the DDPP model, standing for descriptive, diagnostic, predictive, and prescriptive. The first type or maturity level of this DDPP model is descriptive analytics.

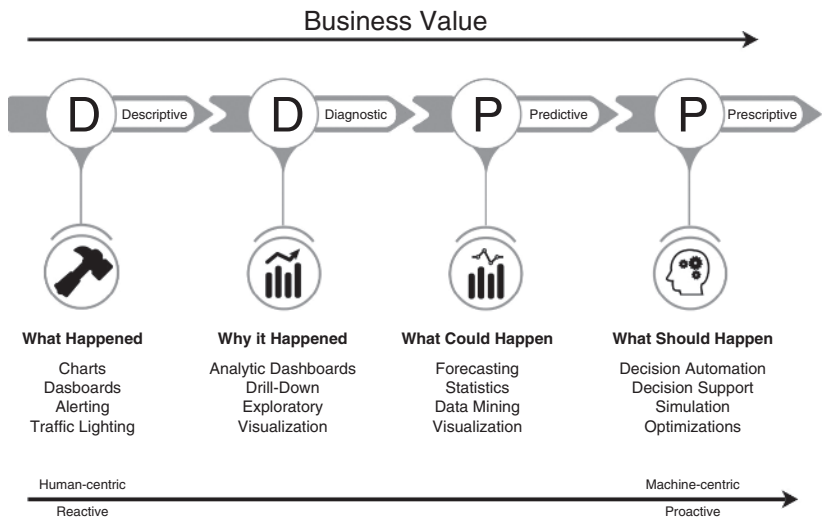


Figure 11 DDPP Model—Types and Maturity of Analytics

This is a reactive level and focuses on the past. In the context of a demand-driven supply chain, this type of analytics focuses on reporting on what has happened. Examples include charts, reports, business intelligence dashboards, alerting, and traffic lighting on key performance indicators (KPIs). Regarding sophistication and business value, this level of maturity and type of analytics provides the lowest benefit of the four levels.

The second level is diagnostic analytics. This type focuses on why something has happened. There are more interrelations between data, and reporting becomes more dynamic and interactive. Examples include interactive business intelligence dashboards with drill-down possibilities. These dashboards are more sophisticated and allow more exploration and visualization of data. The business value shifts to the right, providing more insights from the data.

Predictive analytics is the third level of maturity in the DDPP model. This type of analytics focuses on what could happen and leverages advanced analytics to provide possible outcomes. In the context of demand-driven forecasting, examples include forecasting demand for products based on demand signals. Multiple variables could be included in the analysis to identify possible links, correlations,

and influencing factors, such as tastes, demographics, price, events, weather, location, and so on. Machine learning (ML) and artificial intelligence (AI) could be used to identify the most suitable statistical model to forecast a time series for different products throughout the hierarchy of a product dimension (see Figure 5). Identifying influencing factors and leveraging automated forecast model selection via ML and AI working together is a differentiating benefit of advanced forecasting. Demand signals will vary for different products, and applying the same forecasting models (e.g., autoregressive integrated moving average [ARIMA], exponential smoothing model) across all products will not be as useful as analyzing patterns, causal variables, and trends of the time series, and then applying a suitable forecast model. The computing platform must be intelligent to perform such analysis and have adequate compute (CPU+RAM) resources to complete the task in a time window that supports the business function (e.g., demand planning process cycle). Another example of utilizing ML and AI for demand forecasting could be clustering data from like products (also referred to as surrogate products) to help forecast demand for new products that may share similar traits or attributes. Only a machine could digest such vast amounts of data and spot commonalities and trends that could be applied to forecast products with no historical data. The business value of this level of analytics shifts further to the right.

The final level of maturity in the DDPP model is prescriptive analytics. This type of analytics focuses on providing decision support or automating decisions and providing input into downstream decision and supply chain planning systems. Advanced analytics with machine learning and artificial intelligence could be used to execute simulations and optimizations of possible outcomes and select the most appropriate decision based on data-driven analytics. The sophistication of advanced analysis using the scale and depth of data available, increased automation, and timely decisions all increase the business value to the highest possible in the DDPP model.

Business value through the DDPP types of analytics is further increased by leveraging a hybrid approach of data stores and technologies such as data lakes, databases, data marts, and so on, and then using cloud computing benefits (i.e., elastic scale, automation, ease of management, storage, processing, financial costs, etc.) to ingest

and process hot and cold data in a timely manner. Such integration of components and technologies can lead to complex architecture designs and costs. However, one of the benefits of cloud computing is the ability to leverage specialized supply chain software solutions that are cloud aware (e.g., leverage fundamentals of the cloud computing paradigm). Another benefit is to utilize a platform as a service (PaaS) model in a cloud environment. Such a PaaS model makes it possible for organizations to extend on building blocks of software components and software stacks to address business challenges without having to worry about foundational elements. This could mean an organization could merely spin up a data lake environment or a data warehouse, leverage data ingestion technologies to process hot and cold data, and utilize tools for advanced analytics and visual reporting without having to worry about deployment, management, maintenance, or development of such components.

An example of a software solution stack to help organizations solve demand forecasting challenges through a cloud computing environment is depicted in Figures 12, 13, and 14. This example is based on Microsoft Azure Artificial Intelligence (AI). Azure is the name of the Microsoft cloud, which at the time of writing is one of the top two public cloud vendors in the world. At a high level, the AI services provide the advanced analytics and are the link between the data and presentation layer (see Figure 12). The underlying components required to ingest, store, analyze, and present data (to decision makers or decision systems) are included in the suite. There are multiple choices an organization can make depending on its need

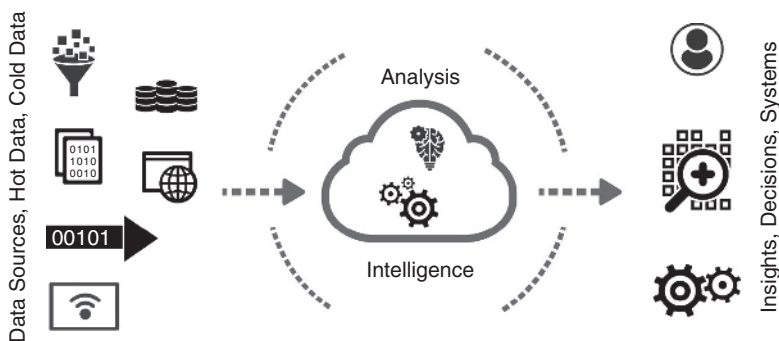


Figure 12 Microsoft AI Example—High Level

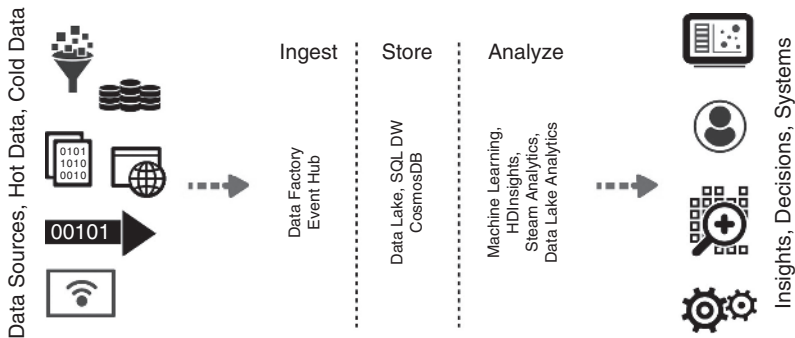


Figure 13 Microsoft AI Services Example

for data (e.g., hot, cold, structured, unstructured, etc.) ingestion and analysis. The design of the Microsoft AI suite has applied principles of the lambda architecture elaborated upon earlier. As depicted in all three diagrams (Figures 12, 13, and 14), data can be from a vast mixture of sources. There is support for hot (fast-moving) and cold (slowly moving) data. Data ingestion is handled by components listed under the Ingest category (see Figure 13). The Microsoft Azure Data Factory is a data ingestion service in the cloud that provides ETL/ELT tasks with automation and scheduling capabilities. Data sources could reside on an organization's premises, already be in the public cloud, or even be from other cloud services (e.g., software as a service [SaaS] applications). Back-end computing nodes are automatically scaled to support the data workloads. The Microsoft Azure Data Factory is a visual design tool used to build a digital data pipeline between data sources and data storage.

Microsoft Azure Event Hubs is also a cloud-based data ingestion service that focuses on hot data. It enables organizations to stream in data and log millions of events per second in near real time. A key use case is that of telemetry data (e.g., IoT devices). It is a cloud-managed service, meaning an organization does not have to worry about development, deployment, maintenance, and the like. These duties are the responsibility of the cloud vendor (in this example it is Microsoft). Event Hubs can be integrated with other cloud services of Microsoft, such as Stream Analytics. The Azure Event Hubs service focuses on near real-time analysis of the hot data being streamed into the cloud, and it can help with rapidly automated decision-making

processes (e.g., anomaly detection). The Microsoft Azure Data Catalog is another example of a cloud service. This service focuses on providing a metadata layer to the disparate sources of data and makes it easier for information/data consumers and data scientists to locate and leverage data. Data remains in its location, and data consumers can utilize tools they are familiar with to access and analyze the data. Data storage in the Microsoft Azure AI services follows principles of polyglot persistence, and an organization can leverage different types of storage types and technologies. Such a hybrid approach provides flexibility and agility to store and later analyze cold and hot data. The decoupling of the data layer makes it possible to store, process, and analyze hot and cold data in parallel utilizing the most suitable data storage technologies for each task. Technologies included in the Microsoft Azure AI services include Data Lake, Microsoft SQL Server Data Warehouse, and a NoSQL database called Cosmos DB. These technologies are cloud services and managed by Microsoft.

The analytics layers (far right of Figure 13) are also cloud services. Included in the Microsoft Azure AI services are open-source technologies (e.g., Hadoop, Spark) and Microsoft technologies such as SQL Server. The back-end architecture scales elastically depending on the demand of the analytics. Management of these components is the responsibility of Microsoft (the public cloud vendor in this example). This makes it easy and financially viable to use these services, as they are all cloud-based. Spinning up and managing such environments on-premises would take days, if not weeks or even months, and would incur substantial up-front investments. Both hot and cold data can be analyzed via this platform. The Microsoft Azure Data Lake Analytics and Azure HDInsight cloud services focus more on cold data and are well suited for investigating unstructured or semistructured data. Microsoft Azure Stream Analytics focuses on hot data. In this context, it receives near real-time data from cloud services such as Microsoft Azure Event Hubs or Microsoft Azure IoT Hubs. Microsoft's Azure Stream Analytics is an example of serverless services as there are no servers or components to manage from an organization's point of view. The public cloud provider (i.e., Microsoft) manages the back end of this service, ensuring it provides the scale required for the analytics of an organization. An organization

that consumes this service pays only for the processing and not for the infrastructure.

Applying such technologies and capabilities of these cloud services to the challenges of demand-driven forecasting can help organizations achieve improved forecast accuracy and business insights in less time, and more cost-effectively. Organizations taking advantage of such possibilities could improve their forecasts and demand planning process by sensing demand from downstream sources and adapting forecasts as new hot data is streamed and analyzed. Demand sensing uses granular downstream sales data (e.g., point of sales data, sales orders) with minimal latency to gain an insight into demand patterns and the impact of demand-shaping programs to refine short-term demand forecasts and inventory positioning to support supply plans of one to six weeks (Chase 2013, 24). An organization could also leverage near real-time data (e.g., online consumer shopping cart) and cold data (e.g., past historical shopping data of consumers) to provide a personalized shopping experience possibly leading to increased sales. Data, analytics, and intelligence are combined to improve demand-driven forecasts or shape demand. Two examples highlight these opportunities and possible benefits.

EXAMPLE

EXAMPLE 1: VENDING MACHINE—IOT DEMAND SENSING

In a simplified view (see Figure 14) in this example, connected (IoT) beverage devices could use telemetry to send near real-time demand signals upstream to demand planning and forecasting systems. Other data sources that are identified as causal factors (i.e., weather in this example) could also be sending hot data back to demand-driven forecast systems. Rolling forecasts could be updated based on newly arriving data. The time horizon of forecasts depends on each organization and the products sold. Perishable products, for example, would require a daily or even hourly forecast, whereas non-perishables (i.e., beverages in this example) could permit a weekly or monthly horizon. Improving the forecasts based on near real-time demand signals could help prevent stock-outs and lost sales. Improved forecast accuracy and timely insights could also assist with goals of preventing high inventory costs (supply exceeding demand).

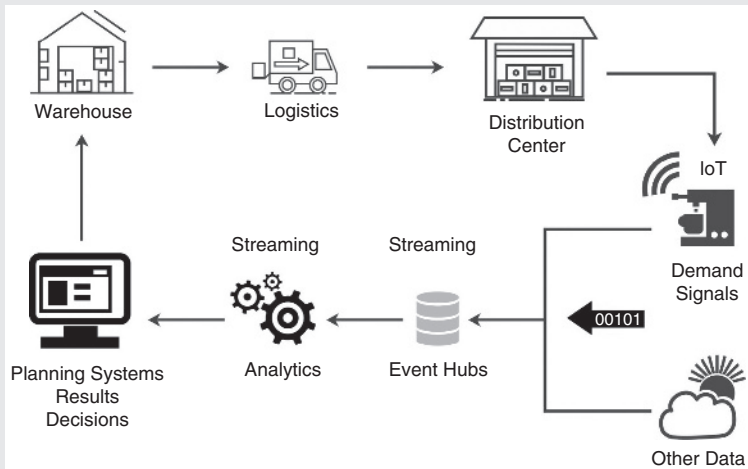


Figure 14 Demand-Driven Forecasting and IoT

EXAMPLE

EXAMPLE 2: ONLINE SHOPPING, DEMAND SHAPING, AND PERSONALIZED RECOMMENDATIONS

In this example, a consumer is shopping online. The surfing data of the current online session is streamed and analyzed in near real time with the help of weblogs. The consumer is a returning buyer, and data from previous purchases is available via cold storage. Advanced analytics is used to identify current shopping behavior and compare it with past actions. Analytics and intelligence (e.g., ML and AI) are used to map the consumer to similar consumers and what those people purchased (e.g., similar gender, age group, hobbies, tastes, etc.). Current inventory stock of the searched-for products as well as other related products is used as an input data variable for analysis.

The machine learning model is used to provide a specifically targeted recommendation to the online shopper. (See Figure 15.) Such a suggestion could take advantage of price sensitivity insights, add sale promotions for related or complementary products, or help reduce inventory overstock by discounting associated products. This provides a personalized shopping experience for consumers and can improve customer satisfaction, increase sale potential, and potentially lower inventory costs. Online stores like Amazon

utilize similar strategies and technologies to help shape demand by using a data-driven approach. New technologies such as bots can take advantage of ML and AI and could provide a simulated human personal shopper experience (e.g., answering questions, providing suggestions) for online consumers. In summary, this example highlights benefits of operating a demand-driven supply chain, sensing demand, and using data and analytics to help shape demand (e.g., enticing consumer to act on purchases via targeted promotions and recommendations).

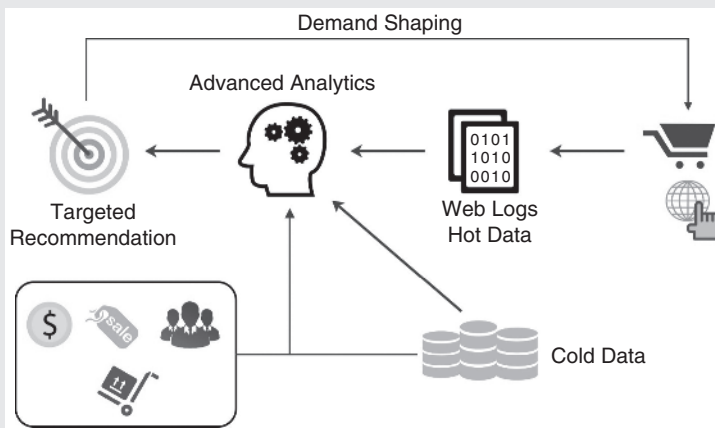


Figure 15 Demand Shaping—Personalized Recommendations

3. COLLABORATION—PEOPLE AND PROCESSES USING A COLLABORATIVE APPROACH

The previously mentioned technologies are powerful and necessary tools for organizations wishing to build and operate a digital supply chain. These technologies should be complemented by a collaboration of demand planners and forecasters. Good practices of forecasting and demand planning should be applied to utilize data insights, yet also leverage domain experience and knowledge. Machine-generated forecasts should not necessarily be overridden by human intervention, as the end result could be decreased forecast accuracy. Small adjustments,

especially frequent ones, reduce forecast accuracy (Fildes et al. 2009, 3–23).

The optimal approach would use a blend between statistical demand-driven forecasts and collaboration between demand planners, as well as supply chain partners. Statistical forecasts can provide valuable decision support, but should not be used to automate the demand planning system completely. Automated decision systems will not be avoidable, but the human factor must remain an independent part of the demand planning process (Spitz 2017, 83). The benefits of improving supply chain management through advanced analytics, IT, process improvements, and becoming a demand-driven supply chain (DDSC) can be helpful to participants in a DDSC chain, as illustrated in Figure 16.

Organizations wishing to optimize their supply chain through digitalization combine disciplines (e.g., data and analytics, demand planners, finance, marketing), and aim to improve their maturity level of being a demand-driven supply chain (DDSC).

DDSCs Have the Potential to Deliver Benefits to All Supply-Chain Participants

	 Raw Material Supplier	 Manufacturer	 Retailer	 Consumer
Reducing inventory	✓	✓	✓	
Decreasing working capital	✓	✓	✓	
Improving forecasting accuracy	✓	✓	✓	
Reducing transportation costs	✓	✓		
Optimizing infrastructure	✓	✓	✓	
Decreasing order-expediting costs	✓	✓	✓	
Reducing other operating costs (such as handling and warehousing)	✓	✓	✓	
Reducing head count (such as planners and buyers)	✓	✓	✓	
Decreasing sales-planning and operations-planning time	✓	✓	✓	
Reducing lost sales		✓	✓	
Improving customer sell-through and satisfaction			✓	✓

✓ = Strong benefit ✓ = Partial benefit

Figure 16 DDSC Benefits All Participants—BCG, 2012

Sources: BCG analysis and case experience, and expert interviews.

There may be many advantages of becoming a DDSC, and some of these benefits of following such a strategy are the following (not an exhaustive list):

- Improved business insights
- Timely insights (information at the right time)
- Increased sales opportunities
- Higher sales revenue
- Lower inventory costs
- Better service levels
- Improved customer satisfaction (e.g., product availability, price)
- Omni-channel demand insights
- Possibly enhanced supply chain agility

There should not be an overdependence on IT and new possibilities such as machine learning. Benefits of leveraging such technologies are enormous, and ML accelerates time to business insights, but there is no magic one-button solution to all (Alexander et al. 2016, 10). Artificial intelligence and machine learning can support automation to a certain degree, and for essential parts of an organization's portfolio (e.g., high value, volatile demand) the statistical forecast can provide valuable decision support. It should not be used to automate the demand planning process completely. As mentioned previously, automated decision systems will not be avoidable, but the human factor must remain an independent part of the process (Spitz 2017, 83). The benefits come to organizations when they apply human knowledge and processes of demand planning (e.g., strategic choices, domain knowledge, and information about events), along with machine capabilities such as ML and AI, to computing an enormous size of data (Alexander et al. 2016, 5–10). Even though it is now possible to collect, store, and analyze vast amounts of data, companies must still formulate a strategy for data management (e.g., what and how to collect, store, analyze, and share). A recent survey of 1,500 companies across Europe, the Middle East, and Africa (EMEA) showed that a midsize company with 500 terabytes of data could be spending roughly US\$1.5 million a year in storage and management costs of nonessential data (Alexander et al. 2016, 6).

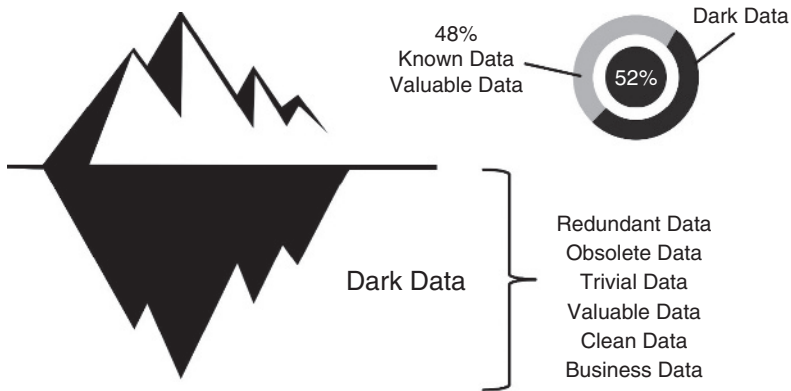


Figure 17 Databerg and Dark Data

Source: Data from Veritas (2016).

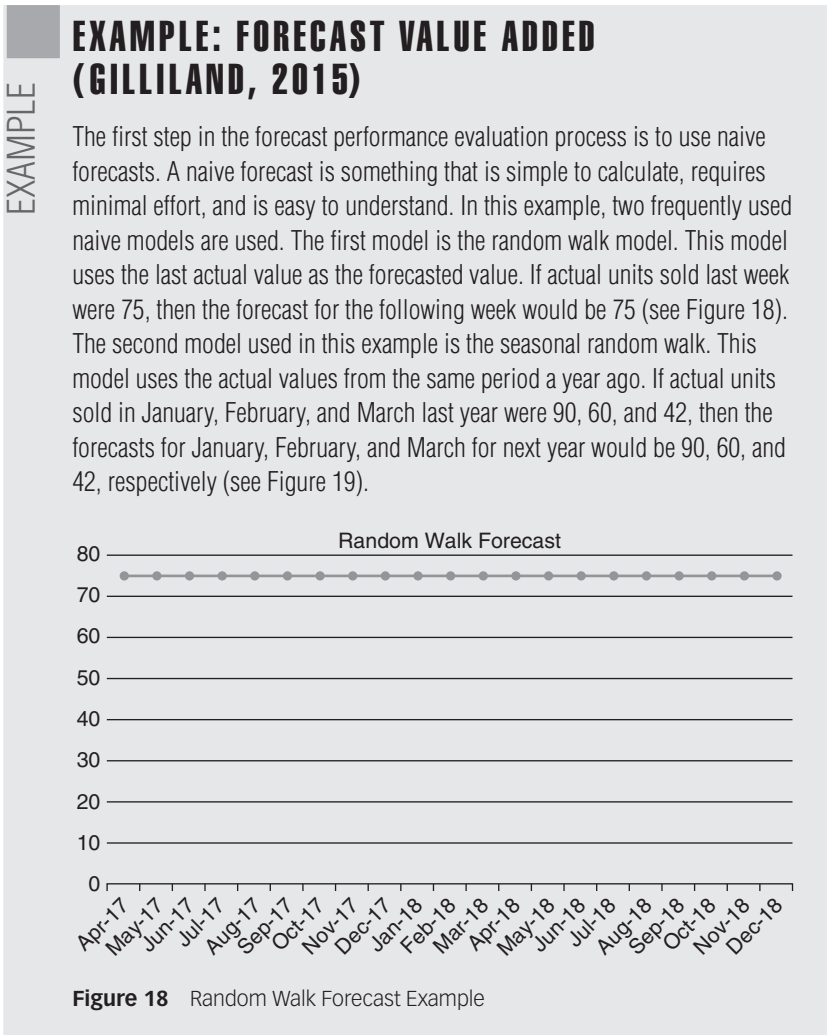
The Veritas Data Genomics Index highlights that over 40 percent of stored data has not been touched and leveraged in more than three years (Veritas 2016, 3). The Veritas Global Databerg Report identifies an average of 52 percent of stored data being “Dark Data,” defined by Veritas as either being redundant, obsolete, or trivial (ROT), or being valuable clean business data (Veritas 2016, 3–5). (See Figure 17.) As there is more and more data, organizations need a good data management strategy to identify useful data, which could then be leveraged in the analytics process and assist with goals of being a demand-driven supply chain. The data value (value to the business, specifically to demand-driven forecasting) should be tested to ensure valuable data is ingested and analyzed. Such tests would probably be repeated over time to make sure an organization is aiming to be as data-smart as it can be.

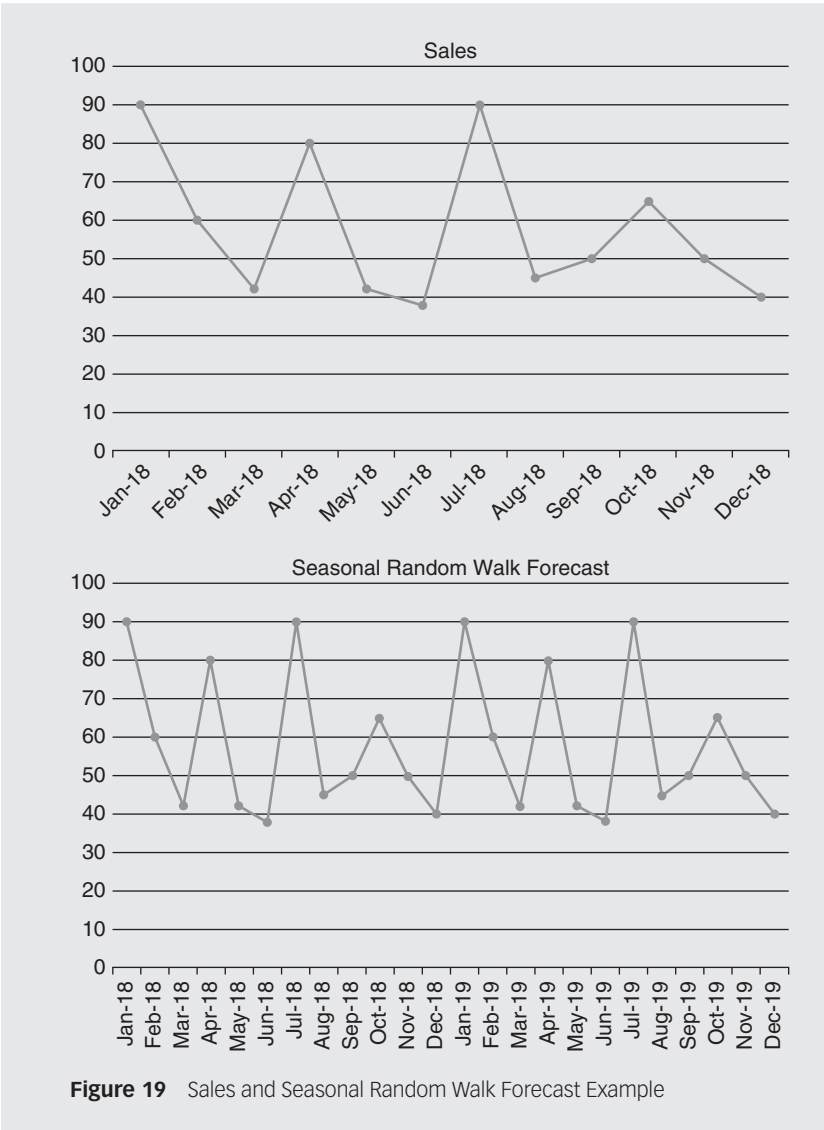
When organizations transition along maturity models of analytics, Big Data, and digitalized supply chain, their needs and capabilities change over time, so it is important to have a process in place that reevaluates data, analytics, and operations. Demand planners have domain knowledge and can have an awareness of events outside of the digital network of data. For example, a purchaser has informed its supplying organization that it will place a large bulk order next month (e.g., lumpy demand), and this information has not yet been captured formally. A purely machine-driven demand forecasting process could be underestimating the actual demand in this example. Another example could be a demand planner knowing about a

planned promotional event that could significantly influence demand. Forecasts could be skewed without such event data, and therefore it is essential to blend data and demand-driven forecasting with human demand planning inputs (achieved via a collaborative approach that seeks the consensus of the machine and the human inputs) to arrive at finalized demand plans.

A survey of forecasters highlighted that 55 percent of the respondents used a mixture of judgment and statistical forecasts, or a judgment-adjusted statistical forecast (Kolassa and Siemsen 2016; Fildes and Petropoulos 2015). Human judgment being widespread in organizational forecasting processes is due to domain-specific knowledge (Lawrence, O'Connor, and Edmundson 2000, 151–160). Domain knowledge and close collaboration across all functions are needed to make the most out of demand sensing and demand shaping (Chase 2013, 24). A two-by-two matrix for forecasting (Croxtton et al. 2002, 51–55) helps to identify the type of forecast needed (e.g., data-driven forecast or people-driven forecast) based on the demand variability and demand volume. Demand Volume (Low and High) are mapped on the *x*-axis while Demand Variability (Low and High) are mapped on the *y*-axis. Low demand variability and low or high demand volume could use data-driven forecasts. High demand variability and low demand volume could use a make-to-order forecast, and high demand variability and high demand volume could use people-driven forecasts (Croxtton et al. 2002, 51–66). Low demand variability and low and high demand volume could use statistical forecasts, high demand variability and low demand volume could use vendor-managed inventory and demand visibility, and high demand variability and high demand volume could use sales and operations planning (S&OP) with collaborative planning and forecast replenishment (Mendes 2011, 42–45). The forecast value added (FVA) metric is a useful tool to help highlight whether or not human judgment and overriding the statistical forecast are adding value. The FVA metric is used to evaluate the performance (positive or negative) of each step and participants in the forecasting process. In this evaluation process, a forecast performance metric is used as the baseline. Deviation from this baseline is then calculated as the forecast value added. If the FVA is positive, then there is a benefit of changing the statistical forecast. If the FVA is negative, then there is no value in altering the statistical forecasts, and it would have been better to leave the

process untouched. The FVA analysis can thus be used to evaluate whether procedures and adjustments made to the statistical forecasts improved or decreased accuracy (Gilliland 2013, 14–18). One of the most common forecast performance metrics used is mean absolute percentage error (MAPE). The lower the MAPE value, the better is the forecast accuracy. Evaluating MAPE and the FVA will highlight when a demand planning process or participant improved the MAPE and added value in overriding a statistical forecast. These concepts are elaborated upon in the following example.





Continuing with the example, the MAPE and FVA metrics are then compared to each other for evaluating the performance of the forecasts and any overrides. The lowest MAPE and highest positive FVA would provide the best business value. Organizations should assess forecasts and processes using a combination of forecast performance

metrics, rather than just relying on one metric. Using FVA, organizations can evaluate whether overriding statistical forecasts added any value.

Organizations must decide on the significance of improvements, and whether using overrides and manual intervention is justified. Review cycles and forecast specialists cost time and money, so it is vital to balance the processes, overrides, and statistical forecasts to provide the most optimal forecasts/demand plans that can be generated in a relatively automated fashion. Such analysis must also bear time pressures in mind, as consumers are becoming more demanding and stress the supply chain in unprecedented ways.

Some software solutions for supply chain optimization include forecasting, inventory optimization, collaborative demand planning, and FVA functionality. The SAS Demand-Driven Planning and Optimization (DDPO) software solution suite is an example of such integration (see Figure 20). A demand signal repository (DSR) collects and consolidates demand signals and other relevant data into function-specific data marts. In this example, the DSR is a blend of the SAS proprietary format (SAS data sets) and a relational database (PostgreSQL). Data is molded into workbench-specific data models (data schemas). The results between business workbenches are shared seamlessly, and this fosters scenarios such as organizations performing demand-driven forecasting with collaborative demand planning, or forecasting with inventory replenishment and optimization, or a combination of all three disciplines (forecasting, collaborative demand planning, and inventory optimization).

An enterprise wishing to reap the most benefits of a demand-driven forecasting strategy should be exploiting advanced technologies in computing resources. Cloud computing provides computing at scale and can do so elastically, scaling out or scaling back in (increasing or decreasing computing power by adding or removing computer servers to process data in parallel). The unlimited computing power, increased agility, on-demand automation, and pay-as-you-go (PAYG) financial model in a cloud makes it very compelling for organizations to leverage such technologies for their demand-driven forecasting needs. Such possibilities and technologies are elaborated upon further in subsequent chapters. Organizations should also leverage capabilities

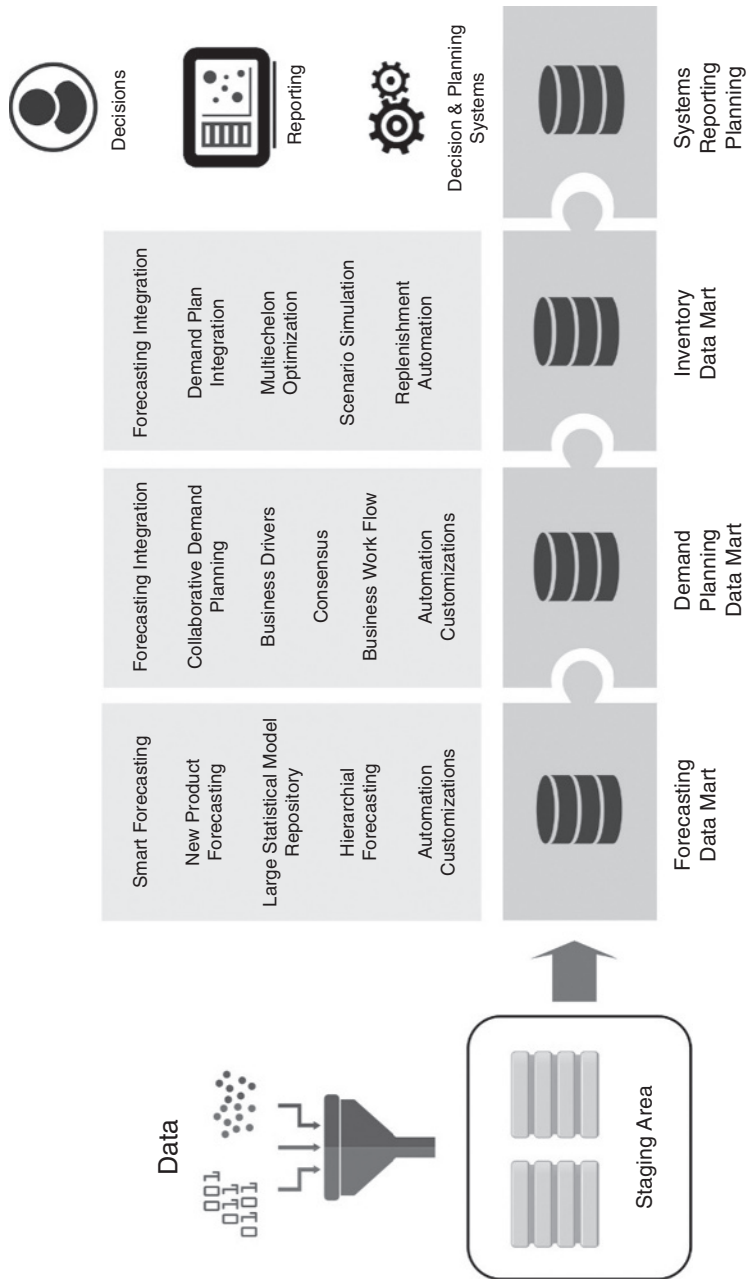


Figure 20 SAS Demand-Driven Planning and Optimization Example

that allow the ingestion and storage of both hot and cold (fast- and slow-moving) data.

Technology advancements, cost reductions, and the increased ease of use now make it possible for artificial intelligence and machine learning to be utilized by the masses. These are all options and possibilities made possible and viable through cloud platforms. One function alone is unlikely to help solve the challenges of a demand-driven supply chain, and it is the combination of all areas that creates exponential value. Last, but not least, these technologies and platforms must be governed by organizational strategies.

Organizations should have a clear data management strategy, and this includes definitions of what to capture, store, analyze, and democratize. Data must be of value, and dark data should be avoided. Equally important are data governance and security. Organizations must be aware of what is being captured, stored, shared, and accessible. New regulations such as the European Union's General Data Protection Regulation (GDPR) must be adhered to, or else organizations can face substantial financial penalties. GDPR came into force in May 2018, and organizations in breach of GDPR could be fined up to 4 percent of annual global turnover, or €20 million (whichever is greater). GDPR does not differentiate between on-premises or a cloud environment, so clear structure, control, governance, and processes are needed. Organizations should select the right balance between statistically derived demand forecasts and collaborative demand planning.

Performance metrics such as FVA should be leveraged to identify possible business benefits of overriding a forecast or altering the demand-driven forecasting processes. Such reviews should be repeated to allow organizations to adapt to changes. Successful organizations keep trying to improve and follow the mantra of excellence being a continuous process and not an accident (BS Reporter 2013, 9). The combination of cloud computing, data, advanced analytics, business intelligence, people, and operations all combined would provide the most value to organizations. Business insights would be available at the right time to make or automate informed and data-driven decisions. The interrelations and benefits

In the context of a cloud-based demand-driven supply chain, the following benefits can be gained by organizations:

Cloud

- A theoretically unlimited pool of computing resources can be utilized.
- Elastically scalable—use more or less computing resources depending on requirements.
- Less to manage—for example, platform as a service (PaaS) and software as a service (SaaS).
- Cost-effective - organizations can gain from economies of scale of cloud providers.

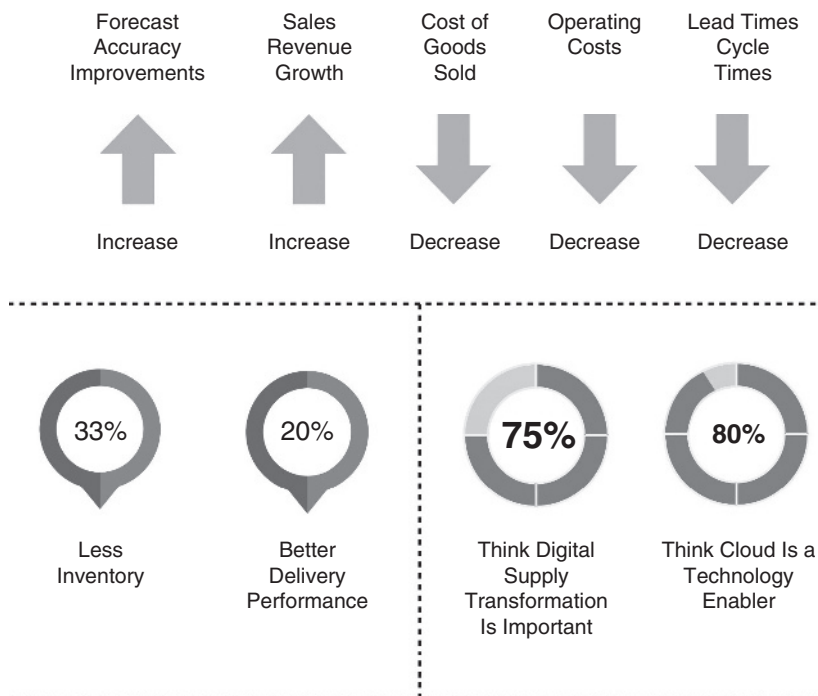


Figure 22 Benefits of Demand-Driven Supply Chain

Sources: (Top) Mendes (2011, 64–65); (bottom left) Budd, Knizek, and Tevelson (2012); (bottom right) Dougados and Felgendreher (2016).

- Pay-as-you-go (PAYG) financial models mean organizations pay for only what they use.
- Cloud data centers are secure, and cloud services are globally accessible.
- Cloud economies make it easier to be highly available and have disaster recovery options.

Data

- Ingest fast- and slow-moving data, and analyze hot or cold data in near real time or in batches.
- Big Data, data lakes, and operational databases.
- Data services scale with organizational needs.

Analytics

- Leverage machine learning and artificial intelligence for advanced analytics.
- Analyze hot data, Big Data, cold data, or a combination of all.
- Use analytics to provide personalized recommendations to downstream consumers.
- Leverage data + AI + ML to create data- and demand-driven insights.

The benefits for an organization digitalizing its supply chain and creating a demand-driven supply chain are illustrated in the infographic shown in Figure 22.

