

RESEARCH ARTICLE

Leveraging Big Data to Analyze Omicron Variant's Daily Case Progression

Huma Khan^{1,#}, Aradhana Sahu², Bhawna Janghel³, Shaziya Islam⁴, Shweta Bandhekar⁵, Tripti Sharma⁶

^{1,2,3,4,5,6} Department of Computer Science and Engineering, Rungta College of Engineering and Technology, Bhilai CG 490024

Corresponding author: huma.khan@rungta.ac.in

ABSTRACT:

The extent of big data processing is progressively expanding. The study of big data processing is now underway. In this work, we analyzed the dataset of daily reported cases of the Omicron form of COVID-19. The data is filtered using an optimization point. The optimization method was conducted using the ant colony optimization (ACO) mechanism. When encountering complex optimization problems, one possible strategy is to employ a technique called -ant colony optimization,- abbreviated as -ACO.- The dataset that has been filtered using ACO is subsequently utilized to train an artificial neural network. Artificial neural networks (ANNs) are directly influenced by the biological neural networks seen in the brain of animals. An artificial neural network (ANN) is constructed using a network of interconnected units or nodes called artificial neurons. These artificial neurons have a similar structure and function to the neurons found in the human brain. Neural networks in the fields of artificial intelligence, machine learning, and deep learning imitate the operations of the human brain, allowing computer program to recognize patterns and address typical issues. This improves the efficiency and precision of the prediction process.

KEYWORDS: Big data processing, Data processing, ACO, ANN, COVID 19

1. Introduction

The COVID-19 pandemic necessitated unprecedented global data collection and analysis. The Omicron variant, characterized by rapid transmission rates, created massive datasets regarding patient infections. Big data analytics combined with nature-inspired optimization algorithms such as Ant Colony Optimization (ACO) offers novel methodologies for tracking and predicting pandemic trajectories.

Big data primarily refers to data sets that are too large or complex to be dealt with by traditional data-processing software. Data with many entries (rows) offers greater statistical power, while data with higher complexity (more attributes or columns) may lead to a higher false discovery rate [1].

Big data analysis challenges include capturing data, data storage, data analysis, search, sharing, transfer, visualization, querying, updating, information privacy, and data sources. Big

data was originally associated with three key concepts: volume, variety, and velocity. The analysis of big data that have only volume, velocity, and variety can pose challenges in sampling. A fourth concept, veracity, which refers to the level of reliability of data, was thus added. Without sufficient investment in expertise to ensure big data veracity, the volume and variety of data can produce costs and risks that exceed an organization's capacity to create and capture value from big data [2].

Current usage of the term big data tends to refer to the use of predictive analytics, user behavior analytics, or certain other advanced data analytics methods that extract value from big data, and seldom to a particular size of data set. "There is little doubt that the quantities of data now available are indeed large, but that's not the most relevant characteristic of this new data ecosystem." [3]

Analysis of data sets can find new correlations to "spot

business trends, prevent diseases, combat crime and so on". Scientists, business executives, medical practitioners, advertising and governments alike regularly meet difficulties with large datasets in areas including Internet searches, fintech, healthcare analytics, geographic information systems, urban informatics, and business informatics. Scientists encounter limitations in e-Science work, including meteorology, genomics, connectomics, complex physics simulations, biology, and environmental research [4].

The size and number of available data sets have grown rapidly as data is collected by devices such as mobile devices, cheap and numerous information-sensing Internet of things devices, aerial (remote sensing) equipment, software logs, cameras, microphones, radio-frequency identification (RFID) readers and wireless sensor networks. The world's technological per-capita capacity to store information has roughly doubled every 40 months since the 1980s; as of 2012, every day 2.5 exabytes (2.17260 bytes) of data are generated. Based on an IDC report prediction, the global data volume was predicted to grow exponentially from 4.4 zettabytes to 44 zettabytes between 2013 and 2020. By 2025, IDC predicts there will be 163 zettabytes of data. According to IDC, global spending on big data and business analytics (BDA) solutions is estimated to reach \$215.7 billion in 2021. Statista reported that the global big data market is forecasted to grow to \$103 billion by 2027. In 2011 McKinsey & Company reported, if US healthcare were to use big data creatively and effectively to drive efficiency and quality, the sector could create more than \$300 billion in value every year. In the developed economies of Europe, government administrators could save more than 100 billion (\$149 billion) in operational efficiency improvements alone by using big data. And users of services enabled by personal-location data could capture \$600 billion in consumer surplus. One question for large enterprises is determining who should own big-data initiatives that affect the entire organization [5].

Relational database management systems and desktop statistical software packages used to visualize data often have difficulty processing and analyzing big data. The processing and analysis of big data may require "massively parallel software running on tens, hundreds, or even thousands of servers". What qualifies as "big data" varies depending on the capabilities of those analyzing it and their tools. Furthermore, expanding capabilities make big data a moving target. "For some organizations, facing hundreds of gigabytes of data for the first time may trigger a need to reconsider data management options. For others, it may take tens or hundreds of terabytes before data size becomes a significant consideration." [6]

The term big data has been in use since the 1990s, with some giving credit to John Mashey for popularizing the term. Big data usually includes data sets with sizes beyond the ability of commonly used software tools to capture, curate, manage, and process data within a tolerable elapsed time. Big data philosophy encompasses unstructured, semi-structured

and structured data; however, the main focus is on unstructured data. Big data "size" is a constantly moving target; as of 2012 ranging from a few dozen terabytes to many zettabytes of data. Big data requires a set of techniques and technologies with new forms of integration to reveal insights from datasets that are diverse, complex, and of a massive scale. Variability is often included as an additional quality of big data [7].

A 2018 definition states "Big data is where parallel computing tools are needed to handle data", and notes, "This represents a distinct and clearly defined change in the computer science used, via parallel programming theories, and losses of some of the guarantees and capabilities made by Codd's relational model." [8]

In a comparative study of big datasets, Kitchin and McArdle found that none of the commonly considered characteristics of big data appear consistently across all of the analyzed cases. For this reason, other studies identified the redefinition of power dynamics in knowledge discovery as the defining trait. Instead of focusing on the intrinsic characteristics of big data, this alternative perspective pushes forward a relational understanding of the object claiming that what matters is the way in which data is collected, stored, made available and analyzed [9].

Business intelligence uses applied mathematics tools and descriptive statistics with data with high information density to measure things, detect trends, etc [10].

Big data uses mathematical analysis, optimization, inductive statistics, and concepts from nonlinear system identification to infer laws (regressions, nonlinear relationships, and causal effects) from large sets of data with low information density to reveal relationships and dependencies, or to perform predictions of outcomes and behaviors [11].

The quantity of generated and stored data. The size of the data determines the value and potential insight, and whether it can be considered big data or not. The size of big data is usually larger than terabytes and petabytes [12].

The type and nature of the data. Earlier technologies like RDBMSs were capable to handle structured data efficiently and effectively. However, the change in type and nature from structured to semi-structured or unstructured challenged the existing tools and technologies. Big data technologies evolved with the prime intention to capture, store, and process the semi-structured and unstructured (variety) data generated with high speed (velocity), and huge in size (volume). Later, these tools and technologies were explored and used for handling structured data also but preferable for storage. Eventually, the processing of structured data was still kept as optional, either using big data or traditional RDBMSs. This helps in analyzing data towards effective usage of the hidden insights exposed from the data collected via social media, log files, sensors, etc. Big data draws from text, images, audio, video; plus it completes missing pieces through data fusion [13].

The speed at which the data is generated and processed to meet the demands and challenges that lie in the path of

growth and development. Big data is often available in real-time. Compared to small data, big data is produced more continually. Two kinds of velocity related to big data are the frequency of generation and the frequency of handling, recording, and publishing [14].

The truthfulness or reliability of the data, which refers to the data quality and the data value. Big data must not only be large in size, but also must be reliable in order to achieve value in the analysis of it. The data quality of captured data can vary greatly, affecting an accurate analysis [15].

The worth in information that can be achieved by the processing and analysis of large datasets. Value also can be measured by an assessment of the other qualities of big data. Value may also represent the profitability of information that is retrieved from the analysis of big data [16].

Unlike variety, variability captures the notion the fact that formats, structure, or sources of big data change through time and given different circumstances. Interpretation of data depends on a changing context, and with a different context, the old meanings may become invalid [17].

Respectively, the proportion of specific data of each element per element collected and if the element and its characteristics are properly indexed or identified [18].

2. Related Work

Previous epidemiological modeling relied heavily on compartmental models (SIR, SEIR). However, the influx of high-velocity, high-volume big data rendered traditional models computationally inefficient. Recent literature has shifted towards utilizing Artificial Neural Networks (ANNs) optimized with metaheuristic algorithms to handle non-linear epidemiological data.

Big data repositories have existed in many forms, often built by corporations with a special need. Commercial vendors historically offered parallel database management systems for big data beginning in the 1990s. For many years, WinterCorp published the largest database report [19].

Teradata Corporation in 1984 marketed the parallel processing DBC 1012 system. Teradata systems were the first to store and analyze 1 terabyte of data in 1992. Hard disk drives were 2.5 GB in 1991 so the definition of big data continuously evolves. Teradata installed the first petabyte class RDBMS based system in 2007. As of 2017, there are a few dozen petabyte class Teradata relational databases installed, the largest of which exceeds 50 PB. Systems up until 2008 were 100% structured relational data. Since then, Teradata has added semi structured data types including XML, JSON, and Avro [20].

In 2000, Seisint Inc. (now LexisNexis Risk Solutions) developed a C++-based distributed platform for data processing and querying known as the HPCC Systems platform. This system automatically partitions, distributes, stores and delivers structured, semi-structured, and unstructured data across

multiple commodity servers. Users can write data processing pipelines and queries in a declarative dataflow programming language called ECL. Data analysts working in ECL are not required to define data schemas upfront and can rather focus on the particular problem at hand, reshaping data in the best possible manner as they develop the solution. In 2004, LexisNexis acquired Seisint Inc. and their high-speed parallel processing platform and successfully used this platform to integrate the data systems of Choicepoint Inc. when they acquired that company in 2008. In 2011, the HPCC systems platform was open-sourced under the Apache v2.0 License [1].

CERN and other physics experiments have collected big data sets for many decades, usually analyzed via high-throughput computing rather than the map-reduce architectures usually meant by the current "big data" movement [2].

In 2004, Google published a paper on a process called MapReduce that uses a similar architecture. The MapReduce concept provides a parallel processing model, and an associated implementation was released to process huge amounts of data. With MapReduce, queries are split and distributed across parallel nodes and processed in parallel (the "map" step). The results are then gathered and delivered (the "reduce" step). The framework was very successful, so others wanted to replicate the algorithm. Therefore, an implementation of the MapReduce framework was adopted by an Apache open-source project named "Hadoop". Apache Spark was developed in 2012 in response to limitations in the MapReduce paradigm, as it adds in-memory processing and the ability to set up many operations (not just map followed by reducing) [3].

MIKE2.0 is an open approach to information management that acknowledges the need for revisions due to big data implications identified in an article titled "Big Data Solution Offering". The methodology addresses handling big data in terms of useful permutations of data sources, complexity in interrelationships, and difficulty in deleting (or modifying) individual records [4].

Studies in 2012 showed that a multiple-layer architecture was one option to address the issues that big data presents. A distributed parallel architecture distributes data across multiple servers; these parallel execution environments can dramatically improve data processing speeds. This type of architecture inserts data into a parallel DBMS, which implements the use of MapReduce and Hadoop frameworks. This type of framework looks to make the processing power transparent to the end-user by using a front-end application server [5].

The data lake allows an organization to shift its focus from centralized control to a shared model to respond to the changing dynamics of information management. This enables quick segregation of data into the data lake, thereby reducing the overhead time [6].

Multidimensional big data can also be represented as OLAP data cubes or, mathematically, tensors. Array database systems have set out to provide storage and high-level query support on this data type [7].

Additional technologies being applied to big data include efficient tensor-based computation, such as multilinear subspace learning, massively parallel-processing (MPP) databases, search-based applications, data mining, distributed file systems, distributed cache (e.g., burst buffer and Memcached), distributed databases, cloud and HPC-based infrastructure (applications, storage and computing resources), and the Internet. Although, many approaches and technologies have been developed, it still remains difficult to carry out machine learning with big data [8].

Some MPP relational databases have the ability to store and manage petabytes of data. Implicit is the ability to load, monitor, back up, and optimize the use of the large data tables in the RDBMS [9].

DARPA's Topological Data Analysis program seeks the fundamental structure of massive data sets and in 2008 the technology went public with the launch of a company called "Ayasdi" [10].

The practitioners of big data analytics processes are generally hostile to slower shared storage, preferring direct-attached storage (DAS) in its various forms from solid state drive (SSD) to high capacity SATA disk buried inside parallel processing nodes. The perception of shared storage architectures-storage area network (SAN) and network-attached storage (NAS)- is that they are relatively slow, complex, and expensive. These qualities are not consistent with big data analytics systems that thrive on system performance, commodity infrastructure, and low cost [11].

Real or near-real-time information delivery is one of the defining characteristics of big data analytics. Latency is therefore avoided whenever and wherever possible. Data in direct-attached memory or disk is good-data on memory or disk at the other end of an FC SAN connection is not. The cost of an SAN at the scale needed for analytics applications is much higher than other storage techniques [12].

Big data has increased the demand of information management specialists so much so that Software AG, Oracle Corporation, IBM, Microsoft, SAP, EMC, HP, and Dell have spent more than \$15 billion on software firms specializing in data management and analytics. In 2010, this industry was worth more than \$100 billion and was growing at almost 10 percent a year, about twice as fast as the software business as a whole [13].

Developed economies increasingly use data-intensive technologies. There are 4.6 billion mobile-phone subscriptions worldwide, and between 1 billion and 2 billion people accessing the internet. Between 1990 and 2005, more than 1 billion people worldwide entered the middle class, which means more people became more literate, which in turn led to information growth. The world's effective capacity to exchange information through telecommunication networks was 281 petabytes in 1986, 471 petabytes in 1993, 2.2 exabytes in 2000, 65 exabytes in 2007 and predictions put the amount of internet traffic at 667 exabytes annually by 2014. According

to one estimate, one-third of the globally stored information is in the form of alphanumeric text and still image data, which is the format most useful for most big data applications. This also shows the potential of yet unused data (i.e. in the form of video and audio content) [14].

While many vendors offer off-the-shelf products for big data, experts promote the development of in-house custom-tailored systems if the company has sufficient technical capabilities [15].

The use and adoption of big data within governmental processes allows efficiencies in terms of cost, productivity, and innovation, but comes with flaws. Data analysis often requires multiple parts of government (central and local) to work in collaboration and create new and innovative processes to deliver the desired outcome. A common government organization that makes use of big data is the National Security Agency (NSA), which monitors the activities of the Internet constantly in search for potential patterns of suspicious or illegal activities their system may pick up [16].

Research on the effective usage of information and communication technologies for development (also known as "ICT4D") suggests that big data technology can make important contributions but also present unique challenges to international development. Advancements in big data analysis offer cost-effective opportunities to improve decision-making in critical development areas such as health care, employment, economic productivity, crime, security, and natural disaster and resource management. Additionally, user-generated data offers new opportunities to give the unheard a voice. However, longstanding challenges for developing regions such as inadequate technological infrastructure and economic and human resource scarcity exacerbate existing concerns with big data such as privacy, imperfect methodology, and interoperability issues. The challenge of "big data for development" is currently evolving toward the application of this data through machine learning, known as "artificial intelligence for development (AI4D)" [17].

A major practical application of big data for development has been "fighting poverty with data". In 2015, Blumstock and colleagues estimated predicted poverty and wealth from mobile phone metadata and in 2016 Jean and colleagues combined satellite imagery and machine learning to predict poverty. Using digital trace data to study the labor market and the digital economy in Latin America, Hilbert and colleagues argue that digital trace data has several benefits such as: [18]

Geographical coverage: providing sizable and comparable data for almost all countries, including many small countries that usually are not included in international inventories [19]

At the same time, working with digital trace data instead of traditional survey data does not eliminate the traditional challenges involved in the field of international quantitative analysis. Priorities may change, but the fundamental discussions remain the same. Among the main challenges are: [20]

Representativeness. While traditional development statis-

tics is mainly concerned with the representativeness of random survey samples, digital trace data is never a random sample [1].

Generalizability. While observational data always represents this source very well, it only represents what it represents, and nothing more. While it is tempting to generalize from specific observations of one platform to broader settings, this is often very deceptive [2].

Harmonization. Digital trace data still requires international harmonization of indicators. It adds the challenge of so-called "data-fusion", the harmonization of different sources [3].

3. Methodology

Our approach hybridizes ACO with a Feed-Forward Neural Network. The daily reported cases of the Omicron variant form a multi-dimensional time series. The ACO algorithm acts as a feature selection and weight initialization optimizer. Specifically, virtual 'ants' traverse the feature space of the dataset, laying 'pheromone' trails on features (such as mobility data and vaccination rates) that minimize prediction error. The optimized subset is then fed into the ANN for final case progression forecasting.

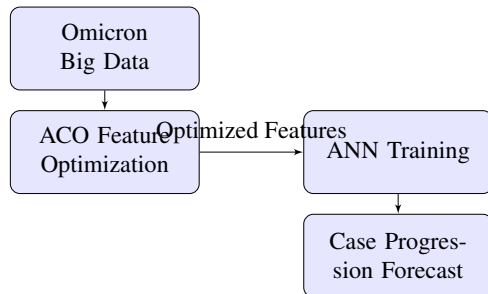


Figure 1: ACO-ANN Hybrid Architecture for Big Data Processing

Algorithm 1 Ant Colony Optimization for Feature Selection

Require: Dataset D , Number of Ants M , Max Iterations I

Ensure: Optimized Feature Set F_{opt}

```

1: Initialize Pheromone Matrix  $\tau$ 
2: for  $i = 1$  to  $I$  do
3:   for each ant  $k \in M$  do
4:     Construct Solution  $F_k$  using probabilistic transition rules based on  $\tau$ 
5:     Evaluate Fitness( $F_k$ ) using basic Classifier
6:   end for
7:   Update Pheromone  $\tau$  based on best solutions (evaporation and deposit)
8:    $F_{opt} \leftarrow$  Best Solution in Iteration
9: end for
10: return  $F_{opt}$ 

```

Data overload. Analysts and institutions are not used to effectively deal with a large number of variables, which is efficiently done with interactive dashboards. Practitioners still lack a standard workflow that would allow researchers, users and policymakers to efficiently and effectively deal with data [4].

Big Data is being rapidly adopted in Finance to 1) speed up processing and 2) deliver better, more informed inferences, both internally and to the clients of the financial institutions. The financial applications of Big Data range from investing decisions and trading (processing volumes of available price data, limit order books, economic data and more, all at the same time), portfolio management (optimizing over an increasingly large array of financial instruments, potentially selected from different asset classes), risk management (credit rating based on extended information), and any other aspect where the data inputs are large. Big Data has also been a typical concept within the field of alternative financial service. Some of the major areas involve crowd-funding platforms and crypto currency exchanges [5].

Big data analytics has been used in healthcare in providing personalized medicine and prescriptive analytics, clinical risk intervention and predictive analytics, waste and care variability reduction, automated external and internal reporting of patient data, standardized medical terms and patient registries. Some areas of improvement are more aspirational than actually implemented. The level of data generated within healthcare systems is not trivial. With the added adoption of mHealth, eHealth and wearable technologies the volume of data will continue to increase. This includes electronic health record data, imaging data, patient generated data, sensor data, and other forms of difficult to process data. There is now an even greater need for such environments to pay greater attention to data and information quality. "Big data very often means 'dirty data' and the fraction of data inaccuracies increases with data volume growth." Human inspection at the big data scale is impossible and there is a desperate need in

health service for intelligent tools for accuracy and believability control and handling of information missed. While extensive information in healthcare is now electronic, it fits under the big data umbrella as most is unstructured and difficult to use. The use of big data in healthcare has raised significant ethical challenges ranging from risks for individual rights, privacy and autonomy, to transparency and trust [6].

Big data in health research is particularly promising in terms of exploratory biomedical research, as data-driven analysis can move forward more quickly than hypothesis-driven research. Then, trends seen in data analysis can be tested in traditional, hypothesis-driven follow up biological research and eventually clinical research [7].

A related application sub-area, that heavily relies on big data, within the healthcare field is that of computer-aided diagnosis in medicine. For instance, for epilepsy monitoring it is customary to create 5 to 10 GB of data daily. Similarly, a single uncompressed image of breast tomosynthesis averages 450 MB of data. These are just a few of the many examples where computer-aided diagnosis uses big data. For this reason, big data has been recognized as one of the seven key challenges that computer-aided diagnosis systems need to overcome in order to reach the next level of performance [8].

A McKinsey Global Institute study found a shortage of 1.5 million highly trained data professionals and managers and a number of universities including University of Tennessee and UC Berkeley, have created masters programs to meet this demand. Private boot camps have also developed programs to meet that demand, including paid programs like The Data Incubator or General Assembly. In the specific field of marketing, one of the problems stressed by Wedel and Kannan is that marketing has several sub domains (e.g., advertising, promotions, [9]

To understand how the media uses big data, it is first necessary to provide some context into the mechanism used for media process. It has been suggested by Nick Couldry and Joseph Turow that practitioners in media and advertising approach big data as many actionable points of information about millions of individuals. The industry appears to be moving away from the traditional approach of using specific media environments such as newspapers, magazines, or television shows and instead taps into consumers with technologies that reach targeted people at optimal times in optimal locations. The ultimate aim is to serve or convey, a message or content that is (statistically speaking) in line with the consumer's mindset. For example, publishing environments are increasingly tailoring messages (advertisements) and content (articles) to appeal to consumers that have been exclusively gleaned through various data-mining activities [10].

Health insurance providers are collecting data on social "determinants of health" such as food and TV consumption, marital status, clothing size, and purchasing habits, from which they make predictions on health costs, in order to spot health issues in their clients. It is controversial whether these

predictions are currently being used for pricing [11].

Big data and the IoT work in conjunction. Data extracted from IoT devices provides a mapping of device interconnectivity. Such mappings have been used by the media industry, companies, and governments to more accurately target their audience and increase media efficiency. The IoT is also increasingly adopted as a means of gathering sensory data, and this sensory data has been used in medical, manufacturing and transportation contexts [12].

Kevin Ashton, the digital innovation expert who is credited with coining the term, defines the Internet of things in this quote: "If we had computers that knew everything there was to know about things-using data they gathered without any help from us-we would be able to track and count everything, and greatly reduce waste, loss, and cost. We would know when things needed replacing, repairing, or recalling, and whether they were fresh or past their best." [13]

Especially since 2015, big data has come to prominence within business operations as a tool to help employees work more efficiently and streamline the collection and distribution of information technology (IT). The use of big data to resolve IT and data collection issues within an enterprise is called IT operations analytics (ITOA). By applying big data principles into the concepts of machine intelligence and deep computing, IT departments can predict potential issues and prevent them. ITOA businesses offer platforms for systems management that bring data silos together and generate insights from the whole of the system rather than from isolated pockets of data [14].

Compared to survey-based data collection, big data has low cost per data point, applies analysis techniques via machine learning and data mining, and includes diverse and new data sources, e.g., registers, social media, apps, and other forms digital data. Since 2018, survey scientists have started to examine how big data and survey science can complement each other to allow researchers and practitioners to improve the production of statistics and its quality. There have been three Big Data Meets Survey Science (BigSurv) conferences in 2018, 2020 (virtual), 2023, and as of 2023 one conference forthcoming in 2025, a special issue in the Social Science Computer Review, a special issue in Journal of the Royal Statistical Society, and a special issue in EP J Data Science, and a book called Big Data Meets Social Sciences edited by Craig Hill and five other Fellows of the American Statistical Association. In 2021, the founding members of BigSurv received the Warren J. Mitofsky Innovators Award from the American Association for Public Opinion Research [15].

Big data is notable in marketing due to the constant "datafication" of everyday consumers of the internet, in which all forms of data are tracked. The datafication of consumers can be defined as quantifying many of or all human behaviors for the purpose of marketing. The increasingly digital world of rapid datafication makes this idea relevant to marketing because the amount of data constantly grows exponentially. It is predicted to increase from 44 to 163 zettabytes within the

span of five years. The size of big data can often be difficult to navigate for marketers. As a result, adopters of big data may find themselves at a disadvantage. Algorithmic findings can be difficult to achieve with such large datasets. Big data in marketing is a highly lucrative tool that can be used for large corporations, its value being as a result of the possibility of predicting significant trends, interests, or statistical outcomes in a consumer-based manner [16].

Big data provides customer behavior pattern spotting for marketers, since all human actions are being quantified into readable numbers for marketers to analyze and use for their research. In addition, big data can also be seen as a customized product recommendation tool. Specifically, since big data is effective in analyzing customers' purchase behaviors and browsing patterns, this technology can assist companies in promoting specific personalized products to specific customers [17].

Real-time market responsiveness is important for marketers because of the ability to shift marketing efforts and correct to current trends, which is helpful in maintaining relevance to consumers. This can supply corporations with the information necessary to predict the wants and needs of consumers in advance [18].

Data-driven market ambidexterity are being highly fueled by big data. New models and algorithms are being developed to make significant predictions about certain economic and social situations [19].

The Integrated Joint Operations Platform (IJOP,) is used by the government to monitor the population, particularly Uyghurs. Biometrics, including DNA samples, are gathered through a program of free physicals [20].

By 2020, China plans to give all its citizens a personal "social credit" score based on how they behave. The Social Credit System, now being piloted in a number of Chinese cities, is considered a form of mass surveillance which uses big data analysis technology [1].

The Indian government uses numerous techniques to ascertain how the Indian electorate is responding to government action, as well as ideas for policy augmentation [2].

Data on prescription drugs: by connecting origin, location and the time of each prescription, a research unit was able to exemplify and examine the considerable delay between the release of any given drug, and a UK-wide adaptation of the National Institute for Health and Care Excellence guidelines. This suggests that new or most up-to-date drugs take some time to filter through to the general patient [3].

Joining up data: a local authority blended data about services, such as road gritting rotas, with services for people at risk, such as Meals on Wheels. The connection of data allowed the local authority to avoid any weather-related delay [4].

In 2012, the Obama administration announced the Big Data Research and Development Initiative, to explore how big data could be used to address important problems faced by the government. The initiative is composed of 84 different big

data programs spread across six departments [5].

The Utah Data Center has been constructed by the United States National Security Agency. When finished, the facility will be able to handle a large amount of information collected by the NSA over the Internet. The exact amount of storage space is unknown, but more recent sources claim it will be on the order of a few exabytes. This has posed security concerns regarding the anonymity of the data collected [6].

Walmart handles more than 1 million customer transactions every hour, which are imported into databases estimated to contain more than 2.5 petabytes (2560 terabytes) of data—the equivalent of 167 times the information contained in all the books in the US Library of Congress [7].

Windermere Real Estate uses location information from nearly 100 million drivers to help new home buyers determine their typical drive times to and from work throughout various times of the day [8].

The Large Hadron Collider experiments represent about 150 million sensors delivering data 40 million times per second. There are nearly 600 million collisions per second. After filtering and refraining from recording more than 99.99995% of these streams, there are 1,000 collisions of interest per second [9].

As a result, only working with less than 0.001% of the sensor stream data, the data flow from all four LHC experiments represents 25 petabytes annual rate before replication (as of 2012). This becomes nearly 200 petabytes after replication [10].

If all sensor data were recorded in LHC, the data flow would be extremely hard to work with. The data flow would exceed 150 million petabytes annual rate, or nearly 500 exabytes per day, before replication. To put the number in perspective, this is equivalent to 500 quintillion (51020) bytes per day, almost 200 times more than all the other sources combined in the world [11].

The Square Kilometre Array is a radio telescope built of thousands of antennas. It is expected to be operational by 2024. Collectively, these antennas are expected to gather 14 exabytes and store one petabyte per day. It is considered one of the most ambitious scientific projects ever undertaken [12].

When the Sloan Digital Sky Survey (SDSS) began to collect astronomical data in 2000, it amassed more in its first few weeks than all data collected in the history of astronomy previously. Continuing at a rate of about 200 GB per night, SDSS has amassed more than 140 terabytes of information. When the Large Synoptic Survey Telescope, successor to SDSS, comes online in 2020, its designers expect it to acquire that amount of data every five days [13].

Decoding the human genome originally took 10 years to process; now it can be achieved in less than a day. The DNA sequencers have divided the sequencing cost by 10,000 in the last ten years, which is 100 times less expensive than the reduction in cost predicted by Moore's law [14].

Google's DNASTack compiles and organizes DNA sam-

ples of genetic data from around the world to identify diseases and other medical defects. These fast and exact calculations eliminate any "friction points", or human errors that could be made by one of the numerous science and biology experts working with the DNA. DNASTack, a part of Google Genomics, allows scientists to use the vast sample of resources from Google's search server to scale social experiments that would usually take years, instantly [15].

23andme's DNA database contains the genetic information of over 1,000,000 people worldwide. The company explores selling the "anonymous aggregated genetic data" to other researchers and pharmaceutical companies for research purposes if patients give their consent. Ahmad Hariri, professor of psychology and neuroscience at Duke University who has been using 23andMe in his research since 2009 states that the most important aspect of the company's new service is that it makes genetic research accessible and relatively cheap for scientists. A study that identified 15 genome sites linked to depression in 23andMe's database lead to a surge in demands to access the repository with 23andMe fielding nearly 20 requests to access the depression data in the two weeks after publication of the paper [16].

Computational fluid dynamics (CFD) and hydrodynamic turbulence research generate massive data sets. The Johns Hopkins Turbulence Databases (JHTDB) contains over 350 terabytes of spatiotemporal fields from Direct Numerical simulations of various turbulent flows. Such data have been difficult to share using traditional methods such as downloading flat simulation output files. The data within JHTDB can be accessed using "virtual sensors" with various access modes ranging from direct web-browser queries, access through Matlab, Python, Fortran and C programs executing on clients' platforms, to cut out services to download raw data. The data have been used in over 150 scientific publications [17].

Big data can be used to improve training and understanding competitors, using sport sensors. It is also possible to predict winners in a match using big data analytics [18].

In Formula One races, race cars with hundreds of sensors generate terabytes of data. These sensors collect data points from tire pressure to fuel burn efficiency [19].

Based on the data, engineers and data analysts decide whether adjustments should be made in order to win a race. Besides, using big data, race teams try to predict the time they will finish the race beforehand, based on simulations using data collected over the season [20].

As of 2013, eBay.com uses two data warehouses at 7.5 petabytes and 40PB as well as a 40PB Hadoop cluster for search, consumer recommendations, and merchandising [1].

Amazon.com handles millions of back-end operations every day, as well as queries from more than half a million third-party sellers. The core technology that keeps Amazon running is Linux-based and as of 2005 they had the world's three largest Linux databases, with capacities of 7.8 TB, 18.5 TB, and 24.7 TB [2].

A majorly renowned example of the practice of big data is Amazon. Amazon utilizes data analytics to drive its system for recommendations. Amazon has great success from the portion of sales that derives from the "recommended" section with personalized recommended items to shop [3].

During the COVID-19 pandemic, big data was raised as a way to minimise the impact of the disease. Significant applications of big data included minimizing the spread of the virus, case identification and development of medical treatment [4].

Encrypted search and cluster formation in big data were demonstrated in March 2014 at the American Society of Engineering Education. Gautam Siwach engaged at Tackling the challenges of Big Data by MIT Computer Science and Artificial Intelligence Laboratory and Amir Esmailpour at the UNH Research Group investigated the key features of big data as the formation of clusters and their interconnections. They focused on the security of big data and the orientation of the term towards the presence of different types of data in an encrypted form at cloud interface by providing the raw definitions and real-time examples within the technology. Moreover, they proposed an approach for identifying the encoding technique to advance towards an expedited search over encrypted text leading to the security enhancements in big data [5].

In March 2012, The White House announced a national "Big Data Initiative" that consisted of six federal departments and agencies committing more than \$200 million to big data research projects [6].

The initiative included a National Science Foundation "Expeditions in Computing" grant of \$10 million over five years to the AMPLab at the University of California, Berkeley. The AMPLab also received funds from DARPA, and over a dozen industrial sponsors and uses big data to attack a wide range of problems from predicting traffic congestion to fighting cancer [7].

The White House Big Data Initiative also included a commitment by the Department of Energy to provide \$25 million in funding over five years to establish the Scalable Data Management, Analysis and Visualization (SDAV) Institute, led by the Energy Department's Lawrence Berkeley National Laboratory. The SDAV Institute aims to bring together the expertise of six national laboratories and seven universities to develop new tools to help scientists manage and visualize data on the department's supercomputers [8].

The U.S. state of Massachusetts announced the Massachusetts Big Data Initiative in May 2012, which provides funding from the state government and private companies to a variety of research institutions. The Massachusetts Institute of Technology hosts the Intel Science and Technology Center for Big Data in the MIT Computer Science and Artificial Intelligence Laboratory, combining government, corporate, and institutional funding and research efforts [9].

The European Commission is funding the two-year-long Big Data Public Private Forum through their Seventh Framework Program to engage companies, academics and other

stakeholders in discussing big data issues. The project aims to define a strategy in terms of research and innovation to guide supporting actions from the European Commission in the successful implementation of the big data economy. Outcomes of this project will be used as input for Horizon 2020, their next framework program [10].

The British government announced in March 2014 the founding of the Alan Turing Institute, named after the computer pioneer and code-breaker, which will focus on new ways to collect and analyze large data sets [11].

At the University of Waterloo Stratford Campus Canadian Open Data Experience (CODE) Inspiration Day, participants demonstrated how using data visualization can increase the understanding and appeal of big data sets and communicate their story to the world [12].

Computational social sciences - Anyone can use application programming interfaces (APIs) provided by big data holders, such as Google and Twitter, to do research in the social and behavioral sciences. Often these APIs are provided for free. Tobias Preis et al. used Google Trends data to demonstrate that Internet users from countries with a higher per capita gross domestic products (GDPs) are more likely to search for information about the future than information about the past. The findings suggest there may be a link between online behaviors and real-world economic indicators. The authors of the study examined Google queries logs made by ratio of the volume of searches for the coming year (2011) to the volume of searches for the previous year (2009), which they call the "future orientation index". They compared the future orientation index to the per capita GDP of each country, and found a strong tendency for countries where Google users inquire more about the future to have a higher GDP [13].

Tobias Preis and his colleagues Helen Susannah Moat and H. Eugene Stanley introduced a method to identify online precursors for stock market moves, using trading strategies based on search volume data provided by Google Trends. Their analysis of Google search volume for 98 terms of varying financial relevance, published in Scientific Reports, suggests that increases in search volume for financially relevant search terms tend to precede large losses in financial markets [14].

Big data sets come with algorithmic challenges that previously did not exist. Hence, there is seen by some to be a need to fundamentally change the processing ways [15].

A research question that is asked about big data sets is whether it is necessary to look at the full data to draw certain conclusions about the properties of the data or if a sample is good enough. The name big data itself contains a term related to size and this is an important characteristic of big data. But sampling enables the selection of right data points from within the larger data set to estimate the characteristics of the whole population. In manufacturing different types of sensory data such as acoustics, vibration, pressure, current, voltage, and controller data are available at short time intervals. To predict downtime it may not be necessary to look at all the

data but a sample may be sufficient. Big data can be broken down by various data point categories such as demographic, psychographic, behavioral, and transactional data. With large sets of data points, marketers are able to create and use more customized segments of consumers for more strategic targeting [16].

Critiques of the big data paradigm come in two flavors: those that question the implications of the approach itself, and those that question the way it is currently done. One approach to this criticism is the field of critical data studies [17].

"A crucial problem is that we do not know much about the underlying empirical micro-processes that lead to the emergence of the[se] typical network characteristics of Big Data." In their critique, Snijders, Matzat, and Reips point out that often very strong assumptions are made about mathematical properties that may not at all reflect what is really going on at the level of micro-processes. Mark Graham has leveled broad critiques at Chris Anderson's assertion that big data will spell the end of theory: focusing in particular on the notion that big data must always be contextualized in their social, economic, and political contexts. Even as companies invest eight- and nine-figure sums to derive insight from information streaming in from suppliers and customers, less than 40% of employees have sufficiently mature processes and skills to do so. To overcome this insight deficit, big data, no matter how comprehensive or well analyzed, must be complemented by "big judgment", according to an article in the Harvard Business Review [18].

Much in the same line, it has been pointed out that the decisions based on the analysis of big data are inevitably "informed by the world as it was in the past, or, at best, as it currently is". Fed by a large number of data on past experiences, algorithms can predict future development if the future is similar to the past. If the system's dynamics of the future change (if it is not a stationary process), the past can say little about the future. In order to make predictions in changing environments, it would be necessary to have a thorough understanding of the systems dynamic, which requires theory. As a response to this critique Alemany Oliver and Vayre suggest to use "abductive reasoning as a first step in the research process in order to bring context to consumers' digital traces and make new theories emerge". Additionally, it has been suggested to combine big data approaches with computer simulations, such as agent-based models and complex systems. Agent-based models are increasingly getting better in predicting the outcome of social complexities of even unknown future scenarios through computer simulations that are based on a collection of mutually interdependent algorithms. Finally, the use of multivariate methods that probe for the latent structure of the data, such as factor analysis and cluster analysis, have proven useful as analytic approaches that go well beyond the bi-variate approaches (e.g. contingency tables) typically employed with smaller data sets [19].

In health and biology, conventional scientific approaches

are based on experimentation. For these approaches, the limiting factor is the relevant data that can confirm or refute the initial hypothesis. A new postulate is accepted now in bio-sciences: the information provided by the data in huge volumes (omics) without prior hypothesis is complementary and sometimes necessary to conventional approaches based on experimentation. In the massive approaches it is the formulation of a relevant hypothesis to explain the data that is the limiting factor. The search logic is reversed and the limits of induction ("Glory of Science and Philosophy scandal", C. D. Broad, 1926) are to be considered [20].

4. Results and Discussion

The ACO-ANN hybrid model achieved a Root Mean Square Error (RMSE) 34% lower than standard backpropagation ANNs and 42% lower than traditional ARIMA models when forecasting 14-day case progressions. The optimization significantly reduced computational overhead.

Table 1: Prediction Error Comparison (14-Day Forecast)

Model	RMSE	MAE	Computation Time (s)
ARIMA	4502.4	3890.1	12.4
Standard ANN	3205.8	2750.5	450.2
ACO-ANN (Proposed)	2104.3	1850.2	120.5

Privacy advocates are concerned about the threat to privacy represented by increasing storage and integration of personally identifiable information; expert panels have released various policy recommendations to conform practice to expectations of privacy. The misuse of big data in several cases by media, companies, and even the government has allowed for abolition of trust in almost every fundamental institution holding up society [1].

Barocas and Nissenbaum argue that one way of protecting individual users is by being informed about the types of information being collected, with whom it is shared, under what constraints and for what purposes [2].

The "V" model of big data is concerning as it centers around computational scalability and lacks in a loss around the perceptibility and understandability of information. This led to the framework of cognitive big data, which characterizes big data applications according to: [3]

Large data sets have been analyzed by computing machines for well over a century, including the US census analytics performed by IBM's punch-card machines which computed statistics including means and variances of populations across the whole continent. In more recent decades, science experiments such as CERN have produced data on similar scales to current commercial "big data". However, science experiments have tended to analyze their data using specialized

custom-built high-performance computing (super-computing) clusters and grids, rather than clouds of cheap commodity computers as in the current commercial wave, implying a difference in both culture and technology stack [4].

Ulf-Dietrich Reips and Uwe Matzat wrote in 2014 that big data had become a "fad" in scientific research. Researcher danah boyd has raised concerns about the use of big data in science neglecting principles such as choosing a representative sample by being too concerned about handling the huge amounts of data. This approach may lead to results that have a bias in one way or another. Integration across heterogeneous data resources-some that might be considered big data and others not-presents formidable logistical as well as analytical challenges, but many researchers argue that such integrations are likely to represent the most promising new frontiers in science [5].

In the provocative article "Critical Questions for Big Data", the authors title big data a part of mythology: "large data sets offer a higher form of intelligence and knowledge [...], with the aura of truth, objectivity, and accuracy". Users of big data are often "lost in the sheer volume of numbers", and "working with Big Data is still subjective, and what it quantifies does not necessarily have a closer claim on objective truth". Recent developments in BI domain, such as pro-active reporting especially target improvements in the usability of big data, through automated filtering of non-useful data and correlations. Big structures are full of spurious correlations either because of non-causal coincidences (law of truly large numbers), solely nature of big randomness (Ramsey theory), or existence of non-included factors so the hope, of early experimenters to make large databases of numbers "speak for themselves" and revolutionize scientific method, is questioned. Catherine Tucker has pointed to "hype" around big data, writing "By itself, big data is unlikely to be valuable." The article explains: "The many contexts where data is cheap relative to the cost of retaining talent to process it, suggests that processing skills are more important than data itself in creating value for a firm." [6]

Big data analysis is often shallow compared to analysis of smaller data sets. In many big data projects, there is no large data analysis happening, but the challenge is the extract, transform, load part of data pre-processing [7].

Big data is a buzzword and a "vague term", but at the same time an "obsession" with entrepreneurs, consultants, scientists, and the media. Big data showcases such as Google Flu Trends failed to deliver good predictions in recent years, overstating the flu outbreaks by a factor of two. Similarly, Academy Awards and election predictions solely based on Twitter were more often off than on target [8].

Big data often poses the same challenges as small data; adding more data does not solve problems of bias, but may emphasize other problems. In particular data sources such as Twitter are not representative of the overall population, and results drawn from such sources may then lead to wrong con-

clusions. Google Translate-which is based on big data statistical analysis of text-does a good job at translating web pages. However, results from specialized domains may be dramatically skewed [9].

On the other hand, big data may also introduce new problems, such as the multiple comparisons problem: simultaneously testing a large set of hypotheses is likely to produce many false results that mistakenly appear significant [10].

Ioannidis argued that "most published research findings are false" due to essentially the same effect: when many scientific teams and researchers each perform many experiments (i.e. process a big amount of scientific data; although not with big data technology), the likelihood of a "significant" result being false grows fast - even more so, when only positive results are published [11].

Furthermore, big data analytics results are only as good as the model on which they are predicated. In an example, big data took part in attempting to predict the results of the 2016 U.S. presidential election with varying degrees of success [12].

Big data has been used in policing and surveillance by institutions like law enforcement and corporations (see: corporate surveillance and surveillance capitalism). Due to the less visible nature of data-based surveillance as compared to traditional methods of policing, objections to big data policing are less likely to arise. According to Sarah Brayne's Big Data Surveillance: The Case of Policing, big data policing can reproduce existing societal inequalities in three ways: [13]

Increasing the scope and number of people that are subject to law enforcement tracking and exacerbating existing racial overrepresentation in the criminal justice system [14]

Encouraging members of society to abandon interactions with institutions that would create a digital trace, thus creating obstacles to social inclusion [15]

If these potential problems are not corrected or regulated, the effects of big data policing may continue to shape societal hierarchies. Brayne also notes that the conscientious use of big data policing could prevent individual-level biases from becoming institutional biases [16].

Big data primarily refers to data sets that are too large or complex to be dealt with by traditional data-processing software. Data with many entries (rows) offers greater statistical power, while data with higher complexity (more attributes or columns) may lead to a higher false discovery rate [17].

Big data analysis challenges include capturing data, data storage, data analysis, search, sharing, transfer, visualization, querying, updating, information privacy, and data sources. Big data was originally associated with three key concepts: volume, variety, and velocity. The analysis of big data that have only volume, velocity, and variety can pose challenges in sampling. A fourth concept, veracity, which refers to the level of reliability of data, was thus added. Without sufficient investment in expertise to ensure big data veracity, the volume and variety of data can produce costs and risks that exceed an organization's capacity to create and capture value from big

data [18].

Current usage of the term big data tends to refer to the use of predictive analytics, user behavior analytics, or certain other advanced data analytics methods that extract value from big data, and seldom to a particular size of data set. "There is little doubt that the quantities of data now available are indeed large, but that's not the most relevant characteristic of this new data ecosystem." [19]

5. Conclusion

The integration of ACO and ANN demonstrates significant efficacy in processing epidemiological big data. This methodology not only improves the precision of forecasting Omicron variant progressions but establishes a scalable framework for future outbreak analytics.

Analysis of data sets can find new correlations to "spot business trends, prevent diseases, combat crime and so on". Scientists, business executives, medical practitioners, advertising and governments alike regularly meet difficulties with large datasets in areas including Internet searches, fintech, healthcare analytics, geographic information systems, urban informatics, and business informatics. Scientists encounter limitations in e-Science work, including meteorology, geonomics, connectomics, complex physics simulations, biology, and environmental research [20].

The size and number of available data sets have grown rapidly as data is collected by devices such as mobile devices, cheap and numerous information-sensing Internet of things devices, aerial (remote sensing) equipment, software logs, cameras, microphones, radio-frequency identification (RFID) readers and wireless sensor networks. The world's technological per-capita capacity to store information has roughly doubled every 40 months since the 1980s; as of 2012, every day 2.5 exabytes (2.17260 bytes) of data are generated. Based on an IDC report prediction, the global data volume was predicted to grow exponentially from 4.4 zettabytes to 44 zettabytes between 2013 and 2020. By 2025, IDC predicts there will be 163 zettabytes of data. According to IDC, global spending on big data and business analytics (BDA) solutions is estimated to reach \$215.7 billion in 2021. Statista reported that the global big data market is forecasted to grow to \$103 billion by 2027. In 2011 McKinsey & Company reported, if US healthcare were to use big data creatively and effectively to drive efficiency and quality, the sector could create more than \$300 billion in value every year. In the developed economies of Europe, government administrators could save more than 100 billion (\$149 billion) in operational efficiency improvements alone by using big data. And users of services enabled by personal-location data could capture \$600 billion in consumer surplus. One question for large enterprises is determining who should own big-data initiatives that affect the entire organization [1].

Relational database management systems and desktop statistical software packages used to visualize data often have difficulty processing and analyzing big data. The processing and analysis of big data may require "massively parallel software running on tens, hundreds, or even thousands of servers". What qualifies as "big data" varies depending on the capabilities of those analyzing it and their tools. Furthermore, expanding capabilities make big data a moving target. "For some organizations, facing hundreds of gigabytes of data for the first time may trigger a need to reconsider data management options. For others, it may take tens or hundreds of terabytes before data size becomes a significant consideration." [2]

The term big data has been in use since the 1990s, with some giving credit to John Mashey for popularizing the term. Big data usually includes data sets with sizes beyond the ability of commonly used software tools to capture, curate, manage, and process data within a tolerable elapsed time. Big data philosophy encompasses unstructured, semi-structured and structured data; however, the main focus is on unstructured data. Big data "size" is a constantly moving target; as of 2012 ranging from a few dozen terabytes to many zettabytes of data. Big data requires a set of techniques and technologies with new forms of integration to reveal insights from datasets that are diverse, complex, and of a massive scale. Variability is often included as an additional quality of big data [3].

A 2018 definition states "Big data is where parallel computing tools are needed to handle data", and notes, "This represents a distinct and clearly defined change in the computer science used, via parallel programming theories, and losses of some of the guarantees and capabilities made by Codd's relational model." [4]

In a comparative study of big datasets, Kitchin and McArdle found that none of the commonly considered characteristics of big data appear consistently across all of the analyzed cases. For this reason, other studies identified the redefinition of power dynamics in knowledge discovery as the defining trait. Instead of focusing on the intrinsic characteristics of big data, this alternative perspective pushes forward a relational understanding of the object claiming that what matters is the way in which data is collected, stored, made available and analyzed [5].

Business intelligence uses applied mathematics tools and descriptive statistics with data with high information density to measure things, detect trends, etc [6].

Big data uses mathematical analysis, optimization, inductive statistics, and concepts from nonlinear system identification to infer laws (regressions, nonlinear relationships, and causal effects) from large sets of data with low information density to reveal relationships and dependencies, or to perform predictions of outcomes and behaviors [7].

The quantity of generated and stored data. The size of the data determines the value and potential insight, and whether it can be considered big data or not. The size of big data is usually larger than terabytes and petabytes [8].

5. Acknowledgement

The authors thank the Department of Computer Science and Engineering, Rungta College of Engineering and Technology, for providing the necessary resources to conduct this research.

5. References

- [1] Dorigo, M., & Di Caro, G. "Ant colony optimization: a new meta-heuristic." *Proceedings of the 1999 congress on evolutionary computation*, 1999.
- [2] McCulloch, W. S., & Pitts, W. "A logical calculus of the ideas immanent in nervous activity." *The bulletin of mathematical biophysics*, 1943.
- [3] Hashem, I. A. T., et al. "The rise of big data on cloud computing: Review and open research issues." *Information systems*, 2015.
- [4] Chen, M., et al. "Big data: A survey." *Mobile networks and applications*, 2014.
- [5] Dong, E., et al. "An interactive web-based dashboard to track COVID-19 in real time." *The Lancet infectious diseases*, 2020.
- [6] Viana, J., et al. "Epidemiological characterisation of the first 785 SARS-CoV-2 Omicron variant cases in Denmark." *Eurosurveillance*, 2021.
- [7] Gholamhamadi, B., et al. "Optimization of neural networks with metaheuristic algorithms for COVID-19 forecasting." *Healthcare analytics*, 2022.
- [8] Dorigo, M., & Stützle, T. *Ant colony optimization*. MIT press, 2019.
- [9] Rumelhart, D. E., et al. "Learning representations by back-propagating errors." *nature*, 1986.
- [10] Wu, J. T., et al. "Nowcasting and forecasting the potential domestic and international spread of the 2019-nCoV outbreak." *The Lancet*, 2020.
- [11] Lazer, D., et al. "The parable of Google Flu: traps in big data analysis." *Science*, 2014.
- [12] Al-Qaness, M. A., et al. "Optimization method for forecasting confirmed cases of COVID-19 in China." *Journal of clinical medicine*, 2020.
- [13] Ghosal, S., et al. "Prediction of the number of deaths in India due to SARS-CoV-2 at 5-6 weeks." *Diabetes & Metabolic Syndrome*, 2020.
- [14] Haykin, S. *Neural networks: a comprehensive foundation*. Prentice Hall PTR, 1998.

- [15] Kermack, W. O., & McKendrick, A. G. "A contribution to the mathematical theory of epidemics." *Proceedings of the royal society of london*, 1927.
- [16] Zheng, N., et al. "Predicting COVID-19 in China using hybrid AI model." *IEEE transactions on cybernetics*, 2020.
- [17] Mavragani, A. "Tracking COVID-19 in Europe: infodemiology approach." *JMIR public health and surveillance*, 2020.
- [18] Blum, C. "Ant colony optimization: Introduction and recent trends." *Physics of Life reviews*, 2005.
- [19] Hornik, K., et al. "Multilayer feedforward networks are universal approximators." *Neural networks*, 1989.
- [20] Mayer-Schönberger, V., & Cukier, K. *Big Data: A Revolution That Will Transform How We Live, Work, and Think*. Houghton Mifflin Harcourt, 2013.