

Chapter 2: Foundations and Futures of Cloud Computing in Distributed Intelligence

2.1. Introduction to Cloud Computing

Cloud computing has become one of the most discussed technologies in recent years and is considered a significant enabler in the bucket of supporting technologies for the inclusion of intelligent applications in all walks of life. Despite the name, “cloud computing” does not imply technology based on computers in clouds but rather distributed intelligence. If computer science is in its infancy today, distributed intelligence is the toddler stage for the long road the discipline has ahead.

The basic idea of cloud computing has existed in the community of computer users since the inception of the Internet; the technology has laid the foundation for intelligent service delivery. Distributed intelligence emerged when the explosion of services hosted in the cloud—the need for executing instructions, and therefore the application logic—was also outsourced from the local device to another flash of intelligence somewhere in the cloud. An explosion in intelligence implies an explosion in volumes of data and distributed applications being submitted relentlessly to the cloud for timely processing to assist in intelligent decision-making. Meeting this explosion demands more—intelligence gathered from the data centres is being sent back and distributed for execution to the fog nodes at the end of the network, close to the edge nodes, and the devices at the bottom to enable decisions quickly and real-time control; this is a nascent concept that is still evolving.

2.1.1. Understanding Cloud Computing: Foundations and Significance

Cloud computing is a general term for any delivered service that is consumed in a distributed environment, where the consumers need not possess knowledge of or control over the technological infrastructure supporting the service. Despite its new-sounding

name, cloud computing is not a breakthrough innovation but rather a natural evolution of distributed intelligence.

Cloud computing consists of three service layers: infrastructure, platform, and software. Each service layer can be offered as a public, private, or hybrid cloud. Cloud computing offers enormous possibilities, but also challenges associated with the infrastructure, applications, and the information itself. The protection of information residing in the cloud is critical yet difficult to achieve because of the lack of physical control over the information. From an information security perspective, cloud computing presents numerous challenges such as protecting confidential information, supporting secure and efficient access control, and maintaining privacy and trust. Services deployed in highly distributed environments are continuously subjected to attacks, thereby jeopardizing the entire infrastructure. Other major concerns stem from a compliance point of view.

2.2. Historical Context of Distributed Intelligence

The vision of distributed artificial intelligence, calling for the construction of intelligent networks around knowledge centers, was suggested in 1984 as a research direction in the field of artificial intelligence. Such a vision has more recently been further identified as a major challenge for future networks in the IEEE Communications Society. Distributed intelligence and knowledge ecosystems rely on coordinated actions of different intelligent and knowledge-oriented nodes communicating through global and suitably designed wired, wireless, and optical networks, so as to deliver a common set of services to the users and communities of interest [1-3]. An appropriate networking architecture was recently described under the generic name of Knowledge Networks Architecture.

The ability to minimize the central control and thus prevent centralized control failure has been a desire since the origins of telecommunications networks. Early circuit-switch networks, such as the standard telephone network, were typically hierarchical, whereas the Internet protocol has developed and is operating as a peer-to-peer networking architecture. This peer-to-peer approach with all distributed nodes at the same level has been very successful in creating the Internet. However, the peer-to-peer approach is not effective in offering differentiated and guaranteed quality of service. The need for differentiation requires a new networking model—one that provides differentiated services while eliminating or very substantially reducing the need for a permanent centralized control entity. Cloud Networks architecture, a multilevel clustering networking model, was proposed to provide such differentiated services in the knowledge ecosystems for distributed intelligence.

2.2.1. Evolution and Milestones in Distributed Intelligence

Cloud computing denotes a model for seamlessly delivering computing resources — servers, storage, networking, software, analytics, and intelligence — over the Internet. These offerings can be rapidly deployed and scaled. By presenting users with attractive and useful services developed by companies that excel at the respective functionalities, cloud computing has significantly influenced and benefited almost all individuals who use digital technologies in their daily lives.

Despite the emerging prominence of cloud computing, efforts to build large-scale distributed systems go back further. Cloud-computing-related concepts have now penetrated into fields other than computing. While cloud computing is often conceived as a way to enhance machine intelligence, it also serves as an external brain for humans by delivering utilities that support decision-making and intelligent behaviors. Since that time, cloud computing has exerted a strong influence on people's general life and their behavior of using digital technologies.



Fig 2 . 1 : Cloud Computing: Delivering Resources and Influencing Daily Life

2.3. Key Concepts in Cloud Computing

Cloud computing has drawn increased attention and received different definitions. According to Buyya, cloud computing “refers to both the applications delivered as services over the Internet and the hardware and system software in the datacentres that provide those services.” Zhang defines cloud computing as a “style of computing in which dynamically scalable and often virtualized resources are provided as a service over the Internet.” NIST defines cloud computing as “a model for enabling convenient, on-demand network access to a shared pool of configurable computing resources that can be rapidly provisioned and released with minimal management effort or service provider interaction.” The core characteristics of different viewpoints can be summarized from the aspects of resource model and service orientation: The resource model perspective considers cloud computing as a large-scale, distributed, parallel, virtualized, and dynamically-scaled resource pool; the service orientation perspective describes cloud computing from the viewpoint of client demands, emphasizing the evaluation of demand-driven services and applications.

The essence of cloud computing embodies two core concepts—Cloud as a Service (CaaS) and Service as a Service (SaaS). Regarding Cloud as a Service, users regard the cloud as a source of resources and services accessed via the Internet, focusing on the dynamic provisioning of resources: when the number of requests increases, abundant resources can be provisioned quickly; when the number of requests decreases, the excess resources are released, thereby maximizing support for clients while minimizing cost. The storage-oriented cloud is a major example. Service as a Service refers to a role-exchange situation with the relationship between clients and service providers. In cases of Service Provider as a Service, the platform provider seeks to take advantage of the capabilities of Infrastructure as a Service providers and create a higher-level service to customers. A typical example of this pattern is Platform as a Service. In Provisioning as a Service, an advanced client serves as a resource allocator, managing and distributing computing tasks to multiple service providers simultaneously [2-4]. For example, an advanced user may adopt several Infrastructure as a Service providers and allocate computing tasks to these providers to achieve the optimal balance for cost and performance.

Agents in Different Clouds

Agents in different clouds are working in a distributed way. Depending on the nature of the job, they can head towards the nearest cloud or the most powerful one, with the final location chosen in such a way as to avoid overloading any cloud and reduce the time of execution. Even though the agent is performing in different clouds, the execution logic remains the same for all the instances active at a particular time.

2.3.1. Infrastructure as a Service (IaaS)

The concept of Infrastructure as a Service (IaaS) was envisioned in 1961 by John McCarthy, who predicted that computer power and communication might one day be sold like commodity services. The operational realization of such a vision appeared decades later with the advent of cloud services. IaaS delves into the single-tenant full- and near-virtualized service offering, followed by the development of multi-tenant virtualized service delivery with tenant-controlled software stacks. Since the launch of Amazon Elastic Compute Cloud (Amazon EC2) in 2006, only a handful of private cloud implementations have emerged.

IaaS provisioning services, including Amazon EC2, GoGrid, and flexiScale; storage services such as Amazon Simple Storage Service S3 (Amazon S3), Nirvanix, and Mosso; and content delivery services like Akamai and Amazon CloudFront remain in wide common use. Amazon Simple Queue Service (SQS), Google's AppEngine Task Queue, and Windows Azure Content Delivery Network represent the growth of task-oriented provisioning services in IaaS public clouds. Within a decade, this cloud service type has become quite commercially mature and is being actively exploited by academic users. IaaS evolved into a set of services: the load distribution that combined compute, storage, and CDN capabilities; the Service-Oriented Infrastructure (SOI) API, which provided middleware services such as read/write queue operations; and the control infrastructure provisioning interface that enabled intelligent management of control levels of implementation of these constituent services in a cloud environment.

2.3.2. Platform as a Service (PaaS)

In the Platform as a Service (PaaS) model, a cloud provider offers computing platforms, including hardware and software tools hosted on its infrastructure. New application development for the infrastructure is simplified with the use of a set of tools built on top of the underlying platform. The user develops new applications without having to worry about platform availability, installation, maintenance, and updates. Usage is typically priced based on execution time and storage requirements, with traditional licensing models applied for specific software components. Users can select from a pool of hardware configurations such as processing capacity, memory size, and storage capacity, similar to IaaS.

PaaS provides reduced risk without investment, simplified issues with legal permission acquisition, and greater flexibility around timing. The PaaS model is hailed as the future for development engineers, offering a free setup environment. The goal is to provide the development environment and tools entirely in the cloud, eliminating concerns with software installation and maintenance. In PaaS, specifically the deployment model, a

provider develops a cloud operating environment. A company establishes its operating environment on this cloud operating environment through the internet. The platform created can then be made accessible to parties involved only in the development process, thus limiting the support environment. This approach facilitates wider user collaboration during the development phase.

2.3.3. Software as a Service (SaaS)

Although the term Software as a Service is often used in connection with the Platform as a Service concept, it can be argued that SaaS is a completely separate concept and model. SaaS is a model in which applications are delivered as services to end-users by means of a cloud infrastructure. No cloud provisioning is required because the SaaS provider controls the applications and interfaces. End-users are unaware of the underlying infrastructure, operating systems, or development platforms used. Examples of this type of application include Web-based e-mail, project management tools, and online retail services.

Although most SaaS users are individual people, SaaS applications can also be used by organizations. The term cloudware is sometimes used to refer to SaaS designed for organizations. It is normally considered to be Web-level software, which is much easier to maintain than standalone desktop applications. Platform as a Service goes one level beyond this to cover programming tools that can be delivered through the browser. It is useful to think of the last two as the business equivalent of e-mail and instant messaging, with Web Services providing the business equivalent of instant messaging. Definitions for cloud computing models and architecture have been proposed, together with the benefits that Cloud provides, mainly for businesses considering executing processes on the cloud.

2.4. Architectural Models of Cloud Computing

Cloud computing is a collective substitute term for “utility computing”, “software plus services”, and “platform computing”. Different fields have different interpretations of the term. It is often viewed as a new style of service-oriented information technology (IT) hosted on the Internet and delivered to the IT users based on AaaS (“as-a-service”) model, e.g., Infrastructure-as-a-Service (IaaS), Platform-as-a-Service (PaaS), and Software-as-a-Service (SaaS). Cloud computing utilizes infrastructures built on data centres, services based on virtual machines and operating systems (OSs), applications, and services offered on the Internet [1,3,5].

The basic architectural models of cloud computing originate from the way the cloud is established and functions. The first architectural model (Figure 2) is based on the similarity to and derivation from utility computing and conceptually consists of three layers—framework, platform, and infrastructure—that build on each other. Framework-level services include a Web 2.0 architecture, Web-oriented call/return protocols (such as Representational State Transfer—REST), and Software as a Service (SaaS) services (including simple, distributed programming models). At the platform level, services include hosting environments and APIs suitable for building the consumer services delivered through the framework. Platform-level services are often used to integrate with the consumer services used in the infrastructure level (at the bottom). Infrastructure-level services include the computer and network resources used to host the consumer services.

2.4.1. Public Cloud

Public cloud services represent one of three main configurations of Cloud resources and services: Public cloud, Private cloud, and Hybrid cloud. A public cloud consists of resources and services publicly available on the Internet. Public cloud providers own and operate the cloud resources and services, like networks, servers, storage facilities, and software applications. These physical assets include datacenters with many hundreds of servers located around the world. Public cloud services are generally available to the broader business community, not just one organization alone.



Fig 2 . 2 : Public Cloud Services

The public cloud provides a range of resources, such as Infrastructure as a Service (IaaS), Platform as a Service (PaaS), and Software as a Service (SaaS). Infrastructure as a Service is an easy way to make resources such as processing, storage, and networking available on-demand to customers, charging them at a monthly or annual subscription rate or on an hourly or daily usage basis. Platform as a Service offers a complete development and deployment environment in the cloud, with resources that enable the provision of everything from simple cloud-based apps to sophisticated, cloud-enabled enterprise applications. Software as a Service[aa27] makes an application available from a cloud service provider to the cloud-consuming organization, which can access it via a web browser or other external interface.

2.4.2. Private Cloud

A private cloud is a cloud infrastructure operated solely for a single organization. This infrastructure may be managed internally or by a third party, and it may be hosted either internally or externally. Unlike public-cloud solutions, a private cloud denies access to the general public and restricts access to a single organization. This means that all hardware and software resources are configured solely for the group's use, making it the most secure cloud architecture compared to its counterparts. Nevertheless, it still offers the usual benefits of cloud services, including scalability, rapid service provisioning, and cost reduction.

Private clouds are particularly attractive for sectors such as finance and government, where data privacy is paramount. Although they require higher capital expenditure due to the need to build and maintain a private cloud infrastructure, they enable organizations to retain higher control over security and data, which are critical corporate concerns [2,4,6]. This type of cloud also serves companies or individuals needing personalized cloud resources for an incidental period. For example, companies in the media and entertainment industry are prime candidates for private clouds, in which their competent teams may work more efficiently.

2.4.3. Hybrid Cloud

Assumptions behind the cloud computing model, state-of-the-art cloud-computing systems, and projects on the horizon are considered. A classical definition of hybrid cloud computing underlines that it combines two or more clouds across a private or specialized cloud and a public cloud. Hybrid cloud systems allow the creation of cloud structures that combine the most suitable parts of private and public clouds, enabling the execution of different parts of a computational task in public or private cloud

components. For example, in the Bank of America model, personal data stays in the private cloud while the analysis of customer interactions occurs in the public cloud.

Assuming the possibility of designing architectures involving more than one cloud using different deployment models or service profiles, the definition is more general, since the service delivery model is not necessarily public for one of the clouds or different for each cloud. In the multi-cloud computing paradigm, components of an application are placed on a number of independent clusters or clouds according to the specific requirements of each component. Such capabilities are for further research, as they are not yet available on the market. Furthermore, concurrency control aspects for hybrid clouds have not received attention. A notable case of hybrid cloud system design and implementation has recently emerged at the University of Concepción.

2.5. Cloud Computing Security Challenges

Following the fast development of cloud computing toward artificial intelligence development, security has become ground zero for the cloud platform's viability. Many cloud providers and users have reported information leakage, data disappearance, and damages to the intellectual property, or service regularity and disaster tolerance, arbitrarily. As the cloud computing architecture has the edge in enhancing business flexibility and agility with the pay-as-you-go payment model, cloud computing suffers from fatal defects in security. They include data confidentiality and privacy, data storage reliability and availability, cross-tenant network security, service availability, fault tolerance, disaster tolerance, and regulatory compliance. Security is a serious problem for both cloud providers and users. For example, a visual imaging sketch consisting of two tigers and buildings was traced to identify the cloud by Google's algorithms that process the over-10-million costs of Google Answer Engine, showing that Google's over-100-million-cost Cloud Engine is not secure or private. This section reviews the cloud computing security challenges and the related significant research efforts to identify potential solutions.

Cloud computing mainly reflects the three characteristics of SAAS, PAAS, and IAAS. The different providers and users form a stack structure with corresponding protection strategies [3,5,7]. IAAS mainly provides the core architecture for virtualization and the distributed network facilities, and the security risks concentrate on the underlying network security, virtualization environment security, etc., of the IAAS layer. The largest number of vulnerabilities was found in the Amazon IAAS platform. Unique Amazon environment vulnerabilities, such as Amazon personal credential caching, Amazon auto-scaling, and Amazon account hijacking, were also found. PAAS systems provide cloud platform services for different API standards in different networks; the security risk lies very much in the potential risk behind the built-in function set in

different platform operation environments. The newly discovered major vulnerabilities in the public platform are that some web services of Microsoft Azure and Google App Engine can be threatened within certain ranges, such as cross-site scripting, SQL injection, cookie stealing, and DoS.

2.5.1. Data Privacy

Privacy is the condition of being free from being observed or disturbed by others. Data privacy safeguards the human right to privacy and ensures that data are collected, further processed, and used in a manner compatible with the data subject's privacy rights. It also guarantees that unauthorized persons cannot access such sensitive data without their managed access. Data privacy is crucial in today's world because vast amounts of personal data are constantly being collected, analyzed, and interlinked.

Despite the inception of data protection when the issue became a source of concerns, the use of data privacy for compliance is still relatively straightforward (such as in GDPR); both obligations on companies and companies' rights about customers—whether the latter conform to a kind of processing or other such states, where GDPR performance does not dictate the level of consumers' performance), regarding both the access companies give to their data and the control (of a similar nature) they exert over it. Data protection still aims to guard people and personal data, with the result that protection is necessary, although not sufficient, to safeguard privacy—even though it is often labeled "data privacy."

2.5.2. Access Control

Access control is central to distributed intelligence and clouds because data and services must be shared with multiple clients and many marked users and machines, all supposedly operating under different permission levels. Consider, for example, data or services in a cloud with the classify attribute. Such data or service cannot be shared with users who do not express the appropriate clear attribute. Otherwise, the necessary attribute comparison and permission decisions must be left to each client, which may not be as trusted, capable, or as consistent as a single service trusted to perform all classification for, say, a government agency.

Distributed AI as a concept has been around for many years. Early proposals for distributing specific AI technologies over different processors date back to the early 1970s. In 1974, Hearsay-I, a speech-understanding program, was distributed over several machines to take advantage of specialized hardware and to run in parallel. The advent of ARPANET, a simple set of networking specifications and protocols, subsequently led to

the distribution of simple expert system technology over a series of machines. ADAGE was an example of a distributed system that integrated a rule-based expert system and a diagnostic drawing program executing on different hosts across the ARPANET.

2.5.3. Compliance Issues

Compliance, along with trust and ethical behavior, is one of the three issues that arise whenever a computer system or a system-of-systems is to be delegated significant authority. Indeed, the vision of cloud computing with serverless functions—at least in part—has long been central to that of the traditional Internet-of-Things, where the devices (things) have limited computational resources and delegate the (central) server much of the computation. A concrete example is a vision of smart traffic lights, where the sensors on the lanes monitor density and periodically submit that information to a server. The server then decides if some lights must change and notifies the relevant ones accordingly. The condition that the server must adhere to can be stated as “never allow traffic in street A to advance toward the intersection while there is still a high density of vehicles in the street crossing it.”

A formal notion of compliance can then be introduced, which guarantees that the certified server will not violate a given traffic condition. This requirement must be checked for every possible sequence of inputs to the server—a classic problem in model-checking—and it can be addressed only by limiting the number of states of the server. Limiting the program’s state space implies a scary trade-off: either the server checks the traffic condition properly and never violates it, or it manages a large state space and can violate the traffic condition.

2.6. Future Trends in Cloud Computing

Cloud computing allows data centers to scale very quickly in order to satisfy increasing amounts of traffic. Closing the gap between the increasing demand and the supply limits of other resources is therefore the major challenge ahead for the field of cloud computing [6-8]. Because data centers are limited by other resources such as bandwidth and power, future developments in cloud computing will need to include new optimization techniques that reduce the use of those resources.

For example, switching costs in the data plane must be minimized to handle the rapid increase in throughput; machine-learning algorithms for traffic prediction and management are likely to be incorporated into the control plane; better design of bandwidth guarantees is needed to provide optimal network utilization; better power

distribution design across the data-center network is needed to reduce electricity consumption.

2.6.1. Edge Computing

Most definitions agree that cloud computing is about avoiding the use of local computing resources, and indeed, Google even instructs developers that an ideal cloud application should never make reference to the local computing environment [Foster et al. (2008)]. Nevertheless, considering the ever-increasing and diverse demands for distributed, heterogeneous, flexible, and scalable computing resources combined with the trend towards the Internet of Things (IoT), a delayed shift from the centralized cloud towards a centralized edge has arisen. This transition is well known in the IoT community as edge computing (Mayer-Schönberger and Cukier [2013]). Figure 6.1 clarifies the relationship between these supporting paradigms and future trends.

Edge computing, also sometimes referred to as fog computing or cloudlet, executes on distributed infrastructure at the edge of the Internet. While its concept is not new, initially it sought to overcome the high latency when exchanging information with the centralized cloud, as well as resilience problems and the centralized cloud acting as a single point of failure. The edge is often mobile and dynamic, which introduces added complexity. The key characteristic of a cloudlet is to be completely localized, so an edge device can also act as a service provider for a different scenario or another user, such as in a disaster or temporary event situation. Unlike other paradigms, the cloudlet has adopted the conventional user–provider view. In contrast to the cloud, the terminal acts as a client, while the cloudlet serves as an infrastructure provider. As a part of the execution environment, the edge cloudlet supports assistant skills such as monitoring or task migration. All the assisted mobile devices depend on services offered by the cloudlet, and that is the reason the latter must also provide special proxy connectors to the application vendors of the mobile terminals.

2.6.2. Quantum Computing Integration

Quantum computers are expected to form the next generation of computing devices. Although current quantum computers are still in their infancy and limited in capabilities, they already demonstrate the enormous potential of quantum computing due to the fact that quantum algorithms can achieve significant speedup for certain problems compared to standard algorithms. One example is Shor’s factoring algorithm, which can solve the factorization problem with polynomial (roughly cubic) complexity. The optimal complexity of the best-known classical algorithm is still subexponential, which is much higher. Shor’s algorithm breaks the security assumptions of almost all classical

encryption methods; therefore, it cannot be run on classical computers. Quantum computers will also accelerate many other algorithms, including search, simulation of quantum mechanical operations, optimization, machine learning, and many more.



Fig 2 . 3 : Quantum Computing: Potential and Challenges

However, quantum computing currently faces a number of hardware- and software-related challenges. For example, a quantum computer with millions of qubits with sufficiently low error rates needs to be constructed to effectively deploy Shor's algorithm. As a result, quantum algorithms will initially be deployed for other solutions that consume far less resources, and a limited number of qubits are available in quantum cloud computing platforms offered by many high-tech companies. Additionally, the software stack (e.g., compiler, error correction, and CPU) for quantum computing systems is not yet mature, which currently lowers the execution efficiency of quantum algorithms. Finally, the lifecycle of digital qubits is tens of milliseconds at maximum, yet the existing queueing delay on quantum cloud computing platforms is often multiple seconds. Some of these issues could be mitigated by integrating quantum computers with distributed clouds. For instance, the near-term quantum processing units (QPUs) are

specialized co-processors possibly utilized in the cloud to improve the execution of certain sub-tasks.

2.6.3. Serverless Architectures

Serverless architectures (also called serverless or Function-as-a-Service—FaaS) break up monolithic applications into smaller functions that perform a single task. These functions execute in response to an event, like putting an item on a queue or changing a state in a database Table. Functions typically are short-lived (a few seconds), stateless, and run in parallel. They are charged by the number of requests and execution duration.

This architecture greatly simplifies running and scaling applications. Instead of managing VMs or containers, developers upload code to a cloud vendor. The cloud provider then scales the application according to demand, handles failures, and administers security. As a result, applications are highly available and developers can focus on their business logic. In addition, event-driven execution lowers costs and improves elasticity, especially when demand is spiky.

2.7. Case Studies of Cloud Computing in Distributed Intelligence

Cloud computing is a recent distributed-computing architecture that can provide the enabling framework for a new distributed intelligence that is central to the IoT–social–emotion convergence. Although cloud computing is quite diverse and continues to evolve, the proposed framework aims at creating intelligence that is distributed relative to a wide range of intelligent decisions and operations. A number of emerging cloud-computing applications that employ distributed intelligence are discussed, followed by the challenges posed and the directions for further research.

Distributed intelligence is not new. Although definitions differ, it generally refers to a situation in which the decisions of individuals—ranging from simple distributed operations to relatively complex decision-making problems—lead to intelligent collective behavior through simultaneous or successive interactions. Distributed intelligence has been a hallmark of social intelligence and/or collective intelligence; it is exemplified by ant colonies, open-source communities, stock markets, and crowdsourced services for data annotation. Ultimately, human society is the most general form of distributed intelligence, constituting intelligence that is distributed across every type of activity and every part of the world. The social networks that mediate these interactions now provide the foundation of a new kind of distributed intelligence that is enabled by distributed computing and emerging technologies associated with the cloud.

2.7.1. Healthcare Applications

Demand for integrated applications via the web has been greater in the health care domain compared to others. Heterogeneity of the applications and the data used makes it necessary to develop new paradigms that utilize the existing networks and databases distributed all over the globe. An integrated application that encapsulates heterogeneous data and applications in a distributed environment is healthcare. Various components such as telemedicine, medical imaging, patient tracking, and medical information management can encapsulate many of these healthcare components at the application level. At the grid level, an encapsulation scheme can hide the heterogeneity of resources, overlay a uniform resource address space over the network, and perform resource scheduling, so that the best-suited resource available can be transparently used.

In a grid environment, not only are different applications handled, but also resources from different fields with different capabilities and constraints can be combined to obtain new knowledge that cannot be obtained from any individual set of resources. For example, in the field of medicine, resources ranging from diagnostic devices and simulation devices to high-performance computing environments and digital libraries for medical information retrieval can be combined to support existing and emerging applications that will revolutionize the face of health care. Such an integrated environment will be helpful for research and training, and even has the potential to improve the health care of people suffering because of a lack of infrastructure in remote parts of the world. Furthermore, the large volume of patient data generated by modern digital patient-care devices has led to the emergence of knowledge discovery techniques and has attracted the attention of researchers, consumers, and manufacturers.

2.7.2. Smart Cities

The Internet of Things (IoT) is progressively developing, with applications aiming to enhance the quality of life for citizens by building smarter and innovative infrastructure and public services. The Internet of services gains significant value through the smart connection of many physical and virtual objects, giving shape to the Internet of Things. IoT-based smart city applications pave the way for intelligent living using communication technologies. Smart cities rely on smart devices that collect large amounts of data on various aspects of life and the environment; however, these devices have limited capacity, and the data generated requires scoring and processing for multiple objectives [8-10]. Many applications and organizations around the world analyze this data and process it, serving multiple objectives and zones. When the data grows to petabytes, the operational process becomes complex. Researchers are therefore trying to build automated smart city applications using Cloud Computing.

The cloud provides a more scalable and efficient platform for integrating different smart city applications. Smart cities use different technologies, and the applications and platforms for these technologies differ. Organizations require large-scale computing resources that can be obtained by utilizing cloud services from third-party providers. Cloud services empower these organizations to outsource data and provide web-based interfaces and services. They charge organizations and users for the resources utilized. City data continues to flood the Internet, generated from smart-meter readings to online social-network feeds. This data requires scalable, cost-effective processing platforms and services in the cloud to harness insights on citizen dynamics and enhance decision-making. Society is changing, and this change is driving the origins of smart cities, which, inherently, are data-driven cities.

2.7.3. Finance Sector Innovations

These features address the essential requirements of cloud computing: reduced capital expenditure through a pay-as-you-go cloud subscription model for storage and computing resources; availability of secure, flexible, and scalable back-end infrastructure and services; remote customer access; and enhanced collaboration and quick time-to-market. For example, banks are going through a change of embracing cloud-based platforms for customer analytics and marketing, mortgage and asset pricing, trade surveillance and compliance, data aggregation and reporting, and risk assessment and Basel II/III compliance. Banks, financial exchanges, asset management firms, and wealth manager firms can store a staggering amount of market, order, and trade data on their private cloud securely and then share that data across organizations. They can use analytics either hosted on the private cloud or hosted by a third-party firm to mine that data for pricing and trade surveillance purposes.

Trading firms in capital markets use private clouds for high-frequency trading, index arbitrage, and pricing complex securities. The money managers can quickly rebalance the portfolio by executing trades on the private cloud with the help of trading algorithms. Asset management firms use private clouds for transfer agency services and settlement services. Big financial exchanges use cloud services for calculating end-of-day settlement prices and margin requirements and performing risk management and Basel-II compliance.

2.8. Conclusion

The paper offers a brief history of distributed artificial and natural intelligence, along with a detailed analysis of the foundations and current state of cloud computing as a new form of distributed computing. On this basis, it provides concrete examples of how cloud

concepts and technologies enrich the current field of distributed artificial intelligence through the Cloud-as-Agent and Agent-as-a-Cloud paradigms. Finally, it projects future developments enabled by forthcoming non-classical computation models.

The paper also synthesizes several approaches to structuring distributed artificial intelligence in a cloud environment, highlighting key directions of evolution. The first approach views clouds as an implementation platform for distributed intelligence, structuring distributed intelligence according to the functional capacities of cloud computing platforms and translating comprehensively developed models and approaches of distributed intelligence into the cloud. The second approach treats distributed intelligence also as a service—a fully functional agent implemented in the cloud. It addresses the technical challenges of collective distributed intelligence, whose members are resources of the cloud. Both approaches define implemented distributed intelligence according to the functional capabilities of the cloud computer. The third approach, by contrast, focuses on how process distribution is implemented in the cloud, regardless of the type of implemented intelligence. In this view, any topic of distributed artificial intelligence, including agent-based intelligence, can be a prospective object of research.



Fig 2 . 4 : The Power of Cloud Computing

2.8.1. Final Reflections and Future Directions in Cloud Computing

The previous overview of cloud computing has presented its foundations, technologies, applications, and future challenges. The starting point of the review was the origins and synthesis of cloud computing, from centralized hardware computers to the aggregation of distributed computing capabilities. A taxonomy of cloud computing services and deployment models was next considered, in terms of types of resources, service models, and scope of accessibility. Particular features and the main enabling technologies of cloud computing were then analysed, and the main applications of cloud computing to business and society were briefly outlined. A SWOT (strengths, weaknesses, opportunities, threats) analysis provided a summary assessment of cloud computing, and the discussion concluded with a future outlook.

Cloud computing represents the use, in a virtualised and location-transparent manner, of large pools of computing resources to provide utility-oriented distributed computing services on demand to external users. Closely related to the concept of distributed artificial intelligence, it differs from previous paradigms by shifting the focus from the centralisation of computing resources of an essential service on the Internet to the centralisation of computing capabilities distributed around the world. This shift helps address a number of drawbacks of the Internet of Services, such as unpredictable performance, scalability limitations, and the high cost of providing elasticity to computing services. Cloud computing offers services at the lowest possible price based on different commercial models, such as subscription or pay-as-you-go. Following the distributed intelligence approach, cloud computing utilises huge amounts of aggregated hardware and software resources made available by different cloud data centres. It does so by managing the spatial and temporal distribution of user requests, while maintaining the quality of the provided services. Future developments will allow computing and knowledge capabilities to be shared and integrated worldwide, not only in computing services, but also by accepting and incorporating the dynamic cooperation of hardware and knowledge resources, making cloud services more intelligent.

References

- [1] Rosendo D, Costan A, Valduriez P, Antoniu G. (2022). Distributed Intelligence on the Edge-to-Cloud Continuum: A Systematic Literature Review. *Journal of Parallel and Distributed Computing*
- [2] Sheelam, G. K. (2025). Agentic AI in 6G: Revolutionizing Intelligent Wireless Systems through Advanced Semiconductor Technologies. *Advances in Consumer Research*.
- [3] Asim M, Wang Y, Wang K, Huang P-Q. (2020). A Review on Computational Intelligence Techniques in Cloud and Edge Computing. *arXiv preprint*.

- [4] Meda, R. (2025). AI-Driven Demand and Supply Forecasting Models for Enhanced Sales Performance Management: A Case Study of a Four-Zone Structure in the United States. *Metallurgical and Materials Engineering*, 1480-1500.
- [5] Tuli S, Mirhakimi F, Pallewatta S, et al. (2022). AI-Augmented Edge and Fog Computing: Trends and Challenges. *arXiv preprint*.
- [6] Inala, R., & Somu, B. (2025). Building Trustworthy Agentic Ai Systems FOR Personalized Banking Experiences. *Metallurgical and Materials Engineering*, 1336-1360.
- [7] Duhok Polytech University team. (2024). Distributed Systems for Machine Learning in Cloud Computing: A Review of Scalable and Efficient Training and Inference. *The Indonesian Journal of Computer Science*, 13(2).
- [8] Kalisetty, S. (2020). Intelligent Supply Chain Ecosystems: Cloud-Native Architectures and Big Data Integration in Retail and Manufacturing Operations. *Open Journal of Educational Research*, 1(1), 1–19.
- [9] Zangana H M, Zeebaree S R M. (2024). Distributed Systems for AI in Cloud Computing: A Review of AI-Powered Applications and Services. *International Journal of Informatics, Information System and Computer Engineering*, 5(1), 11–30.
- [10] Sriram, H. K., Challa, K., & Gadi, A. L. (2025). AI and Cloud-Driven Transformation in Finance, Insurance, and the Automotive Ecosystem: A Multi-Sectoral Framework for Credit Risk, Mobility Services, and Consumer Protection. Anil Lokesh and singreddy, Sneha, *AI and Cloud-Driven Transformation in Finance, Insurance, and the Automotive Ecosystem: A Multi-Sectoral Framework for Credit Risk, Mobility Services, and Consumer Protection* (March 15, 2025).