

Original Article

Hybrid Cloud-Edge Infrastructures for Scalable IIoT AI Deployments

Jennifer Clark

Senior Data Scientist, Innovate Analytics, USA.

Received: 01-03-2020

Revised: 01-17-2020

Accepted: 01-30-2020

Published: 02-03-2020

ABSTRACT

Industrial Internet of Things (IIoT) ecosystems are increasingly reliant on artificial intelligence (AI) to enable predictive analytics, real-time control, and intelligent automation. However, the centralized nature of traditional cloud computing introduces latency, bandwidth, and privacy constraints that limit the real-time applicability of AI models in industrial settings. This paper explores hybrid cloud-edge infrastructures as a scalable solution for deploying AI in IIoT environments. We present an architectural framework that balances compute-intensive model training in the cloud with low-latency inference at the edge. Key challenges including model orchestration, data management, and security in distributed environments are analyzed. Through case studies and performance evaluations, we demonstrate how hybrid architectures can effectively support scalable and resilient AI deployments for a range of industrial applications. Our findings highlight open research challenges and provide recommendations for building robust hybrid IIoT systems.

KEYWORDS

Industrial Internet Of Things (IIoT), Edge Computing, Cloud Computing, Hybrid Architecture, Artificial Intelligence (AI), Model Orchestration, Real-Time Analytics, Scalability, Distributed Systems, Predictive Maintenance.

1. INTRODUCTION

1.1. Background on IIoT (Industrial Internet of Things)

The Industrial Internet of Things (IIoT) represents a transformative paradigm in modern industrial systems by interconnecting machines, sensors, actuators, and software platforms through advanced networking and data analytics. Unlike traditional industrial automation systems, which are often isolated and purpose-specific, IIoT enables the seamless integration of physical assets with digital intelligence, creating cyber-physical systems that are capable of self-monitoring, coordination, and optimization. By leveraging ubiquitous connectivity, IIoT facilitates real-time data acquisition and remote control over production assets, allowing industries such as manufacturing, energy, logistics, and transportation to achieve greater operational visibility, predictive capabilities, and efficiency. These networks generate massive volumes of heterogeneous data, which must be effectively analyzed and acted upon to unlock their full potential. However, the nature of IIoT environments distributed, latency-sensitive, and resource-constrained presents unique architectural and computational challenges, particularly when integrating intelligent analytics.

1.2. The Growing Role of AI in IIoT Environments

Artificial Intelligence (AI), particularly machine learning (ML) and deep learning (DL), has become an integral part of IIoT ecosystems due to its ability to uncover patterns, make predictions, and automate decision-making processes based on real-time data. Applications of AI in IIoT include predictive maintenance, anomaly detection, quality control, demand forecasting, and autonomous system control. These intelligent capabilities not only enhance productivity but also enable systems to adapt to dynamic conditions without human intervention. For instance, machine learning models trained on historical sensor data can predict equipment failures before they occur, reducing downtime and maintenance costs. However, effective AI deployment in IIoT is not trivial; it requires the ability to manage data pipelines, compute-intensive workloads, and model inference across diverse and often constrained devices. As such, integrating AI into IIoT architectures necessitates new design approaches that are responsive to the specific needs of industrial systems—particularly with regard to computational distribution, latency minimization, and resilience.

1.3. Challenges in Scalability, Latency, and Resource Management

Despite the benefits of AI in IIoT, scaling these intelligent systems across industrial sites introduces substantial challenges. One major concern is latency: many industrial applications, such as robotic control or safety monitoring, require real-time or near-real-time responses that cannot tolerate the delays introduced by cloud-based computation. Scalability also poses a challenge, as IIoT networks can consist of thousands of nodes distributed across geographically dispersed locations, each generating data that must be processed, stored, and analyzed. Additionally, resource limitations at the edge—such as limited CPU power, memory, and battery life—complicate the deployment of AI models, particularly those requiring high computational overhead. Managing these resources efficiently while ensuring high availability, low energy consumption, and quality of service remains a central obstacle in building scalable, intelligent IIoT systems.

1.4. Motivation for Hybrid Cloud-Edge Architectures

In response to these limitations, hybrid cloud-edge architectures have emerged as a promising solution to enable scalable, low-latency, and resource-efficient AI deployments in IIoT settings. This architectural model leverages the strengths of both cloud and edge computing: while the cloud provides virtually unlimited computational power and centralized data management for training and updating AI models, the edge enables rapid inference and local decision-making close to the source of data. By intelligently distributing tasks—such as performing real-time inference at the edge while reserving complex analytics and model retraining for the cloud—the hybrid approach offers a balanced trade-off between performance and scalability. Moreover, it enhances system resilience, supports privacy by reducing the need to transmit sensitive data to the cloud, and reduces network bandwidth consumption. The motivation for this paper stems from the growing need to formalize and optimize such architectures for AI-driven IIoT applications across diverse industrial use cases.

1.5. Paper Contributions and Structure

This paper contributes to the growing body of research on intelligent industrial systems by proposing a structured, scalable, and performance-aware hybrid cloud-edge architecture for deploying AI in IIoT environments. The contributions of the paper are threefold: first, we present a detailed architectural framework that combines edge, fog, and cloud layers to support dynamic and distributed AI workloads. Second, we analyze the challenges of orchestrating AI services across heterogeneous environments, with a focus on latency, scalability, and model lifecycle management. Third, we evaluate the architecture through real-world use cases and identify open research directions for future exploration. The remainder of this paper is organized as follows: Section 2 reviews related work in cloud and edge-based IIoT systems; Section 3 details the proposed hybrid architecture; Section 4 explores AI deployment strategies within this architecture; Section 5 discusses scalability and performance optimization; Section 6 addresses security and privacy; Section 7 presents use cases; and Section 8 concludes with future research opportunities.

2. RELATED WORK

2.1. Overview of Current Cloud-Centric and Edge-Centric IIoT Architectures

The literature on IIoT architectures reveals two dominant paradigms: cloud-centric and edge-centric approaches. Cloud-centric IIoT systems rely heavily on centralized data centers for storage, processing, and analytics. In such systems, data generated by IIoT devices is transmitted to the cloud where AI models are trained and executed. While this allows for powerful analytics and cross-device correlation, it suffers from high latency, limited real-time capabilities, and significant bandwidth consumption. On the other hand, edge-centric architectures attempt to address these limitations by shifting computation closer to the data source. Here, AI models are deployed on edge devices such as industrial gateways or embedded systems, allowing for rapid response times and lower data transmission costs. However, this decentralization often leads to challenges in managing distributed resources, ensuring consistency, and handling the computational constraints of edge devices. Both paradigms have merits and drawbacks, motivating the exploration of hybrid approaches that combine their strengths.

2.2. Studies on AI/ML Deployment in IIoT

A growing number of studies have investigated the integration of AI and machine learning techniques within IIoT environments. These efforts range from predictive maintenance systems using time-series analysis to computer vision applications in quality inspection, and reinforcement learning for adaptive process control. Research has shown that embedding AI at the edge can dramatically improve response times and reduce reliance on cloud infrastructure. Other studies have explored cloud-based model training and edge-based inference as a viable trade-off. For example, federated learning has gained traction as a method to train models collaboratively across edge devices without requiring centralized data collection. Despite these advancements, existing research often addresses specific verticals or use cases, and there is a lack of unified frameworks or scalable architectural patterns that holistically manage AI deployment across cloud and edge.

2.3. Limitations of Existing Approaches

While both cloud-centric and edge-centric systems offer partial solutions, they often fall short in real-world industrial settings that demand scalability, robustness, and adaptability. Cloud-only solutions are hindered by network latency, data privacy concerns, and limited autonomy at the edge. Conversely, edge-only deployments struggle with resource limitations, lack of centralized coordination, and difficulties in updating or maintaining AI models across many distributed nodes. Moreover, the absence of standardized orchestration mechanisms and unified data management frameworks across layers leads to fragmented and inefficient systems. These limitations reveal the necessity of a hybrid architectural model that can dynamically allocate resources and adapt to operational constraints while still supporting the complex lifecycle of AI models.

2.4. Comparative Summary

In summary, the current landscape of IIoT architectures and AI deployment strategies is fragmented, with isolated approaches that lack the flexibility and scalability required by modern industrial systems. Cloud-centric models offer computational power but introduce latency and dependency issues; edge-centric models address latency but are limited by hardware capabilities. Recent hybrid approaches attempt to unify these extremes but often lack comprehensive frameworks for orchestration, security, and dynamic workload management. This paper positions itself within this emerging space by proposing a scalable and structured hybrid cloud-edge architecture that builds upon the strengths of existing work while addressing its gaps. The subsequent sections outline how this architecture is designed, implemented, and evaluated to support scalable AI in IIoT.

3. ARCHITECTURE OVERVIEW

3.1. Definition of Hybrid Cloud-Edge Architecture

A hybrid cloud-edge architecture is a distributed computing framework that integrates the centralized power of cloud data centers with the low-latency, on-site responsiveness of edge computing. In the context of Industrial Internet of Things (IIoT), this architecture is designed to support real-time data analytics and intelligent decision-making by enabling AI workloads to be deployed flexibly across different tiers—edge, fog, and cloud—based on computational needs and

latency requirements. The hybrid model mitigates the limitations of purely cloud-based systems, such as bandwidth constraints and high latency, and complements the limited processing capabilities of edge devices. This structure allows industries to dynamically assign workloads—processing data near the source when immediacy is crucial and leveraging cloud resources when scale or storage is needed. It provides a balance of responsiveness, computational efficiency, and system scalability, tailored for the demands of modern industrial environments.

3.2. Key Components

3.2.1. Edge Nodes (Sensors, Gateways, Local Compute)

Edge nodes represent the first point of contact for data in an IIoT system. These include physical sensors that measure parameters like temperature, pressure, vibration, and humidity, along with gateways and embedded devices that offer localized computing capabilities. Often deployed in harsh and resource-constrained environments, edge nodes are responsible for collecting and preprocessing data before it is transmitted to upstream systems. Increasingly, edge devices are capable of executing AI inference tasks, such as detecting anomalies or initiating control signals, thereby reducing latency and network load. Gateways may also aggregate data from multiple sensors and handle protocol translation or basic filtering. Their proximity to data sources makes them indispensable for latency-sensitive applications where real-time decision-making is critical to operations, such as fault detection or safety monitoring.

3.2.2. Fog Nodes

Fog computing introduces an intermediary layer between edge devices and the cloud, offering more robust computing power than the edge but still retaining geographic and logical closeness to the data source. Fog nodes are typically deployed on-site in facilities, substations, or localized data centers and are responsible for managing multiple edge devices while performing mid-level processing and analysis. They can handle tasks such as batch analytics, device coordination, edge-to-cloud data normalization, and intermediate AI inference for larger models. Fog nodes also often include storage capabilities for buffering or archiving local data before it is uploaded to the cloud. This layer improves system reliability, supports load balancing, and helps ensure continuity of operations even during network disruptions.

3.2.3. Cloud Backend

The cloud backend serves as the high-capacity computational core of the architecture. It is responsible for training complex AI models, aggregating large datasets from distributed sources, and executing advanced analytics that require high memory and processing power. The cloud layer also handles global coordination, such as multi-site data fusion, model versioning, long-term storage, and system-wide orchestration. With access to scalable infrastructure and high availability, cloud environments like AWS, Azure, or Google Cloud can support continuous integration/continuous deployment (CI/CD) pipelines for AI models, enabling updates to be pushed to edge and fog nodes as needed. The cloud also plays a key role in visualization and reporting, providing enterprise-wide dashboards and insights that are critical for strategic decision-making.

3.2.4. Communication and Orchestration Layers

The communication and orchestration layers ensure seamless integration and interaction among the architecture's components. Communication protocols such as MQTT, AMQP, CoAP, and OPC UA facilitate lightweight, reliable, and often real-time data exchange between edge, fog, and cloud systems. High-speed networks, including 5G and industrial Ethernet, are vital for latency-sensitive applications. The orchestration layer is responsible for deploying and managing AI workloads across the hierarchy, leveraging container orchestration platforms like Kubernetes for the cloud and lightweight alternatives like K3s or KubeEdge at the edge. These orchestration tools ensure dynamic resource allocation, container health monitoring, service discovery, and automated failover. Together, these layers enable dynamic scaling, continuous service availability, and flexible workload distribution throughout the system.

3.2.5. Deployment Models for AI Workloads

AI workloads in a hybrid architecture are strategically deployed based on latency, bandwidth, and computational requirements. Inference models that require instant response times are deployed on edge nodes, while more computationally intensive model training tasks are handled in the cloud. Fog nodes may host compressed or lightweight versions of these models to perform intermediate tasks like aggregating edge outputs or orchestrating multi-node predictions. The architecture must also support federated and distributed learning models, where training is done collaboratively across multiple nodes with centralized or federated coordination. Deployment pipelines need to include mechanisms for continuous model updates, rollback, and testing across multiple environments, ensuring consistent model behavior regardless of location.

4. AI IN HYBRID CLOUD-EDGE IIOT

4.1. Types of AI Tasks in IIoT: Anomaly Detection, Predictive Maintenance, Optimization

IIoT applications commonly utilize AI for tasks such as anomaly detection, predictive maintenance, and system optimization. Anomaly detection algorithms monitor real-time sensor data to identify deviations from established patterns, signaling potential malfunctions or abnormal conditions. Predictive maintenance models use historical and current operational data to forecast equipment failures, enabling maintenance teams to intervene before costly breakdowns occur. Optimization tasks involve adjusting system parameters to achieve objectives like energy efficiency, throughput maximization, or cost reduction. These AI functions are tightly coupled with the physical environment and thus must be deployed in ways that match the performance and latency requirements of their specific use case.

4.2. Placement Strategies for AI Models: Inference at the Edge, Training in the Cloud

The placement of AI models in a hybrid architecture follows a tiered strategy. Inference models are often deployed at the edge to allow real-time processing of sensor inputs, minimizing latency and improving responsiveness. These models are typically lightweight, optimized for fast execution on embedded hardware. In contrast, model training, retraining, and hyperparameter tuning are computationally intensive and best suited for the cloud, where ample resources and large

datasets are available. This separation of training and inference also facilitates better version control and continuous learning, as updated models can be trained in the cloud and seamlessly pushed to edge nodes via orchestration tools.

4.3. Model Lifecycle Management Across Cloud and Edge

Managing AI models across cloud and edge environments requires a robust lifecycle management strategy that covers development, deployment, monitoring, and updates. In the cloud, models are developed, validated, and trained using large datasets. Once a model is ready for deployment, it is packaged (e.g., as a Docker container or ONNX model) and distributed to edge or fog devices. Monitoring tools at the edge track model performance and detect drift or degradation in accuracy. When a new model version is available, it is deployed using CI/CD pipelines. Lifecycle management must also ensure compatibility with diverse hardware and software stacks across edge nodes, requiring extensive testing and version control.

4.4. Data Flow and Preprocessing Strategies

Efficient data flow is crucial for the performance of AI models in a hybrid architecture. At the edge, raw data from sensors is often preprocessed to remove noise, normalize values, and convert data formats. This preprocessing reduces the amount of data that needs to be transmitted to higher layers and improves inference accuracy. Fog nodes may further aggregate or fuse data from multiple edge devices before sending it to the cloud for archival or training. Data flow strategies must ensure reliability, low latency, and bandwidth efficiency. Intelligent buffering and batching mechanisms help manage variable data rates, while data tagging and labeling facilitate effective learning and analysis in the cloud.

5. SCALABILITY AND PERFORMANCE CONSIDERATIONS

5.1. Resource Management and Orchestration (Kubernetes, K3s, etc.)

Scalability in hybrid IIoT architectures depends heavily on effective resource management. Kubernetes has emerged as a powerful orchestration platform for cloud environments, offering automated deployment, scaling, and health monitoring of containerized AI applications. For edge computing, lightweight variants like K3s or MicroK8s are used due to their minimal resource footprint. These tools allow for decentralized orchestration, ensuring that edge workloads can be dynamically adjusted based on device health, resource availability, or workload intensity. Orchestrators also manage container lifecycle operations like rolling updates, crash recovery, and load scaling, which are essential for sustaining high performance in large-scale deployments.

5.2. Dynamic Workload Distribution

Workload distribution in hybrid systems must be dynamic and context-aware. Rather than statically assigning tasks, the system should evaluate network latency, CPU/GPU availability, memory constraints, and workload priority to determine the optimal placement for each task. For example, a factory may shift quality inspection inference from an overloaded gateway to a nearby fog node during peak hours. Dynamic distribution ensures that the system remains responsive and

avoids bottlenecks by adaptively reallocating tasks based on real-time conditions, making full use of available compute resources while maintaining service-level objectives.

5.3. Latency-Sensitive Decision Making

Latency-sensitive operations, such as safety shutdowns or real-time equipment control, require immediate processing and cannot tolerate the delays associated with cloud round-trips. In such cases, edge deployment of AI models is essential. This ensures deterministic response times and uninterrupted operation, even during network outages. Fog nodes can also play a role by coordinating responses across multiple edge devices in microseconds, especially in systems with collaborative or federated AI. Understanding and minimizing end-to-end latency – from sensor input to AI inference to actuator response – is critical for ensuring operational reliability.

5.4. Load Balancing and Failover Mechanisms

Load balancing and failover are essential mechanisms for maintaining the reliability and availability of distributed systems. Load balancing ensures that incoming data and processing requests are evenly distributed across available resources, preventing bottlenecks and resource exhaustion. Failover mechanisms detect node or service failures and automatically redirect tasks to backup systems or nodes, minimizing disruption. These functions must be built into the orchestration and communication layers and are especially important in critical infrastructure environments where downtime is unacceptable. Health checks, stateful redundancy, and geo-replication techniques are often used to implement robust load balancing and failover.

6. SECURITY AND PRIVACY

6.1. Data Governance in Hybrid Environments

Data governance in a hybrid cloud-edge architecture involves enforcing policies that ensure the integrity, confidentiality, and compliance of data across all tiers. Since IIoT environments involve sensitive operational data, robust governance frameworks must define who has access to what data, how data is stored and transmitted, and how compliance with standards like GDPR or ISO/IEC 27001 is maintained. Edge and fog nodes must adhere to these governance rules even when operating offline or independently from cloud control. Metadata tagging, audit trails, and policy-driven access control help maintain governance across distributed environments.

6.2. Secure Model and Data Transmission

Transmission of data and models between cloud, fog, and edge layers must be secured using encryption and authentication protocols to protect against interception, tampering, or unauthorized access. Technologies like TLS/SSL, VPN tunneling, and zero-trust networking are commonly applied to ensure secure communications. In AI deployments, models themselves must be encrypted or signed to prevent unauthorized manipulation or reverse engineering. Secure update mechanisms also ensure that only authenticated and verified model versions are deployed to field devices, reducing the risk of adversarial attacks or model corruption.

6.3. Trust and Access Control in Distributed AI

Building trust in distributed AI systems requires comprehensive identity and access management (IAM). Each component—be it a sensor, gateway, or cloud service—must authenticate itself and prove its integrity before participating in data exchange or model execution. Role-based access control (RBAC) and attribute-based access control (ABAC) models define what actions users and devices can perform. Blockchain-based identity verification, TPMs (Trusted Platform Modules), and remote attestation can be used to validate device integrity in hostile or untrusted environments. Ensuring trust in model sources, data origin, and computational nodes is crucial for secure and transparent AI deployment in IIoT.

7. CASE STUDIES/USE CASES

7.1. Smart Manufacturing

In smart manufacturing, hybrid cloud-edge architectures facilitate real-time monitoring and predictive maintenance, enhancing operational efficiency. For instance, a capacitor manufacturer faced challenges in identifying root causes of production issues and monitoring data in real-time. By implementing FogHorn's Complex Event Processing (CEP) engine and EdgeML machine learning software, the company analyzed sensor data at the edge, providing real-time analytics and failure alerts to operational technology (OT) staff. This approach empowered the manufacturer to reduce maintenance costs by detecting problems earlier and only calling in technicians when necessary, leading to increased yield and reduced scrap work.

7.2. Predictive Maintenance

Predictive maintenance leverages AI to forecast equipment failures before they occur, minimizing downtime and maintenance costs. A steel manufacturer installed vibration sensors to identify bottlenecks in the production process. However, the data exceeded available bandwidth for transmission to the cloud. By utilizing FogHorn's CEP engine and EdgeML, the manufacturer conducted Fourier analysis at the edge, reducing data sent to the cloud by 90%. This enhancement improved production key performance indicators (KPIs) and laid the foundation for output optimization.

7.3. Energy Management in Industrial Facilities

Energy management in industrial facilities benefits from hybrid architectures by enabling real-time monitoring and optimization of energy consumption. For example, a global oil and gas company monitored industrial pumps with limited communications and computing resources. By deploying FogHorn's CEP engine on existing control gateways, the company analyzed pressure, vibration, and acoustic sensor data at the edge, generating early alerts for abnormal conditions. This proactive approach reduced pump lifecycle costs, cavitation, and equipment damage, transitioning from regularly scheduled maintenance to predictive maintenance.

7.4. Real-World Example of Hybrid Deployment

A notable real-world example of hybrid deployment is LG Electronics' smart factory initiatives. The company integrated Digital Twin technology, big data, and generative AI to improve quality, industrial safety, and equipment management. Through strategic partnerships, including a collaboration with LS ELECTRIC, LG developed the MAVIN-Cloud platform powered by Google Cloud's Anthos technology. This platform facilitates AI-driven quality inspections across manufacturing lines, reducing AI model development time by over 50% and decreasing false defect judgments by 95%.

8. CHALLENGES AND OPEN RESEARCH ISSUES

8.1. Interoperability and Standardization

Interoperability and standardization remain significant challenges in hybrid cloud-edge IIoT architectures. The diverse range of devices, protocols, and platforms used in industrial environments complicates seamless integration. The lack of universal standards for communication protocols, data formats, and device interfaces hinders the efficient exchange of information between heterogeneous systems. Overcoming these challenges requires the development and adoption of open standards and frameworks that promote compatibility and facilitate smooth interoperability across different vendors' systems.

8.2. Real-Time Model Updating and Versioning

Real-time model updating and versioning are critical for maintaining the accuracy and relevance of AI models deployed in hybrid architectures. As industrial processes evolve, AI models must be continuously updated to reflect new data and changing conditions. Implementing real-time model updates involves challenges such as ensuring backward compatibility, managing model dependencies, and minimizing downtime during updates. Effective versioning strategies are essential to track model changes, facilitate rollback procedures, and maintain consistency across edge and cloud environments.

8.3. Scalability of Federated Learning in Hybrid Settings

Federated learning enables collaborative model training across distributed devices without sharing raw data, preserving privacy. However, scaling federated learning in hybrid cloud-edge settings presents challenges related to communication overhead, model convergence, and resource constraints at the edge. Efficient aggregation of model updates, handling heterogeneous data distributions, and ensuring synchronization across devices are critical areas of research. Addressing these challenges will enhance the feasibility and effectiveness of federated learning in large-scale industrial applications.

8.4. Cost Optimization

Cost optimization is a key consideration in deploying hybrid cloud-edge architectures for IIoT applications. While cloud resources offer scalability and advanced analytics capabilities, they also incur operational costs. Edge computing can reduce data transmission costs and latency but may

require significant upfront investment in hardware and infrastructure. Balancing the costs associated with edge and cloud components, including hardware, software, maintenance, and energy consumption, is essential to ensure the economic viability of hybrid deployments.

9. CONCLUSION AND FUTURE WORK

9.1. Summary of Findings

This paper has explored the integration of hybrid cloud-edge architectures in Industrial Internet of Things (IIoT) environments, highlighting their role in enhancing scalability, reducing latency, and optimizing resource management. Through case studies in smart manufacturing, predictive maintenance, and energy management, the effectiveness of these architectures in real-world applications has been demonstrated. The challenges of interoperability, real-time model updating, federated learning scalability, and cost optimization have been identified as critical areas requiring further research and development.

9.2. Potential Directions for Future Research

Future research in hybrid cloud-edge IIoT architectures should focus on developing standardized communication protocols and data formats to improve interoperability across diverse systems. Advancements in real-time model updating mechanisms, including efficient versioning and deployment strategies, are essential to maintain the accuracy and relevance of AI models. Exploring scalable federated learning techniques that address the unique challenges of hybrid environments will enhance collaborative model training without compromising data privacy. Additionally, innovative approaches to cost optimization, such as energy-efficient computing and dynamic resource allocation, will contribute to the economic sustainability of hybrid deployments. Collaborative efforts between academia, industry, and standardization bodies are crucial to drive innovation and address the evolving needs of IIoT applications.

REFERENCES

- [1] Kaur, K., Garg, S., Aujla, G. S., Kumar, N., Rodrigues, J. J. P. C., & Guizani, M. (2018). Edge computing in the industrial Internet of Things environment: Software-defined-networks-based edge-cloud interplay. *IEEE Communications Magazine*, 56(2), 44-51. <https://doi.org/10.1109/MCOM.2018.1700622>
- [2] Aazam, M., Zeadally, S., & Harras, K. A. (2018). Deploying fog computing in industrial Internet of Things and Industry 4.0. *IEEE Transactions on Industrial Informatics*, 14(10), 4674-4682. <https://doi.org/10.1109/TII.2018.2855198>
- [3] Javed, A., Heljanko, K., Buda, A., & Främling, K. (2018). CEFIoT: A fault-tolerant IoT architecture for edge and cloud. In *Proceedings of the IEEE 4th World Forum on Internet of Things (WF-IoT)* (pp. 813-818). IEEE. <https://doi.org/10.1109/WF-IoT.2018.8355149>
- [4] Omoniwa, B., Hussain, R., Javed, M. A., Bouk, S. H., & Malik, S. A. (2018). Fog/Edge Computing-Based IoT (FECIoT): Architecture, Applications, and Research Issues. *IEEE Internet of Things Journal*. <https://doi.org/10.1109/JIOT.2018.2875544>
- [5] Bangui, H., Rakrak, S., Raghay, S., & Buhnova, B. (2018). Moving to the Edge-Cloud-of-Things: Recent Advances and Future Research Directions. *Electronics*, 7(11), 309. <https://doi.org/10.3390/electronics7110309>
- [6] Mahmood, Z., et al. (2018). Toward integrated Cloud-Fog networks for efficient IoT provisioning: Key challenges and solutions. *Future Generation Computer Systems*, 88, 606-613. <https://doi.org/10.1016/j.future.2018.05.015>
- [7] Pizzolli, D., Cossu, G., Santoro, D., Capra, L., Dupont, C., Charalampos, D., De Pellegrini, F., Antonelli, F., & Cretti, S. (2018). Cloud4IoT: A heterogeneous, distributed and autonomic cloud platform for the IoT. *arXiv preprint arXiv:1810.01839*.
- [8] Grover, J., & Garimella, R. M. (2018). Reliable and Fault-Tolerant IoT-Edge Architecture. In *Proceedings of IEEE Sensors 2018*. IEEE. <https://doi.org/10.1109/ICSENS.2018.8589624>.