

Original Article

Time-Series Deep Learning for Proactive Server Reliability Management

Dr. Grace Ndlovu¹, Samuel Johnson²

¹Associate Professor University of Pretoria, South Africa.

²Senior Operations Manager MTN Group, South Africa.

Received: 01-03-2022

Revised: 01-21-2022

Accepted: 01-27-2022

Published: 02-03-2022

ABSTRACT

Modern server systems generate massive volumes of time-series operational data including metrics such as CPU utilization, memory usage, I/O throughput, and network latency. Traditional threshold-based monitoring techniques struggle to capture latent temporal dependencies and evolving performance trends, which often leads to delayed detection of anomalies and reactive maintenance strategies. This paper presents a deep learning-based time-series prediction framework designed to proactively anticipate server reliability issues before they escalate into critical failures. We leverage advanced architectures such as Long Short-Term Memory (LSTM), Temporal Convolutional Networks (TCN), and Transformer encoders to model complex temporal patterns from multivariate telemetry streams. Our approach incorporates adaptive retraining, feature engineering, and explainability methods to balance predictive accuracy with operational interpretability. Experimental results on real-world server telemetry datasets demonstrate significant improvements in early fault detection, reducing mean time to detection (MTTD) and enabling automated preemptive actions. We further discuss integration strategies within DevOps pipelines to transition from reactive monitoring to proactive reliability management.

KEYWORDS

Time-Series Forecasting, Deep Learning, Server Reliability Management, Predictive Maintenance, Anomaly Detection, Lstm (Long Short-Term Memory), Temporal Convolutional, Networks, Transformer Models, Telemetry Analytics.

1. INTRODUCTION

1.1. Background and Motivation

Modern software systems increasingly rely on large-scale server infrastructures deployed across data centers, cloud platforms, and edge environments. These servers support mission-critical applications such as online services, financial systems, healthcare platforms, and enterprise software, where even brief downtime can result in significant economic loss and user dissatisfaction. To ensure continuous service availability, server reliability management has become a central concern in software engineering and site reliability engineering (SRE). Servers continuously emit telemetry data in the form of time-series metrics, including CPU utilization, memory consumption, disk I/O, network throughput, and system latency. This data provides valuable insight into the operational health of systems. However, the sheer volume, velocity, and complexity of such data make manual analysis infeasible. Recent advances in deep learning, particularly time-series modeling, offer an opportunity to automatically learn complex temporal patterns from telemetry data and enable intelligent reliability management. This paper is motivated by the need to leverage deep learning techniques to transform raw telemetry streams into actionable predictions that can proactively prevent failures rather than merely reacting to them.

1.2. Challenges in Traditional Server Reliability Monitoring

Traditional server reliability monitoring techniques primarily rely on static threshold-based rules and reactive alerting mechanisms. In these systems, alerts are triggered when predefined metric thresholds are exceeded, such as CPU usage surpassing a fixed percentage or memory utilization reaching a critical level. While simple to implement, such approaches fail to capture complex temporal dependencies, nonlinear behaviors, and gradual performance degradation. Thresholds that are too strict often generate false positives, leading to alert fatigue, whereas overly lenient thresholds may delay the detection of critical failures. Additionally, traditional monitoring tools operate in isolation, analyzing metrics independently rather than modeling correlations across multiple system components. These limitations make conventional monitoring ineffective in modern distributed and microservices-based architectures, where failures often emerge from subtle interactions between components over time.

1.3. Need for Proactive and Predictive Methods

Given the limitations of reactive monitoring, there is a growing need for proactive and predictive reliability management strategies. Proactive methods aim to anticipate failures before they occur by identifying early warning signals embedded in telemetry data. Predictive models can forecast future system behavior, enabling administrators to take preventive actions such as load balancing, resource scaling, or component replacement. Time-series deep learning models are particularly well suited for this task, as they can learn long-term dependencies, seasonality, and evolving patterns from historical data. By shifting from reactive alerting to predictive analytics, organizations can reduce downtime, improve system resilience, and optimize operational costs. This paradigm shift aligns with modern DevOps and SRE practices that emphasize automation, intelligence, and continuous improvement.

1.4. Paper Contributions

This paper contributes to the field of software reliability engineering by presenting a comprehensive deep learning-based framework for proactive server reliability management using time-series telemetry data. The study systematically formulates the reliability prediction problem, analyzes key server performance indicators, and evaluates multiple deep learning architectures, including LSTM, Temporal Convolutional Networks, and Transformer-based models. The proposed methodology integrates feature engineering, anomaly detection, and explainable AI techniques to ensure both predictive accuracy and interpretability. Extensive experimental analysis demonstrates the effectiveness of deep learning models in early failure detection compared to traditional baseline approaches. Furthermore, the paper discusses practical deployment considerations, highlighting how predictive models can be integrated into real-world monitoring pipelines.

1.5. Paper Organization

The remainder of this paper is organized as follows. Section 2 reviews related work in time-series forecasting, deep learning-based anomaly detection, and existing reliability analytics frameworks. Section 3 formulates the server reliability prediction problem and defines key performance indicators and objectives. Section 4 describes the proposed methodology, including data preprocessing, feature engineering, deep learning architectures, training strategies, anomaly detection mechanisms, and explainability techniques. Subsequent sections present experimental results, deployment considerations, and conclusions.

2. LITERATURE REVIEW

2.1. Time-Series Forecasting in IT Operations

Time-series forecasting has long been used in IT operations to model system performance trends and predict future behavior. Early approaches relied on statistical models such as autoregressive integrated moving average (ARIMA) and exponential smoothing methods, which assume linear relationships and stationary data. While effective for simple workloads, these methods struggle with high-dimensional, non-stationary telemetry data commonly observed in modern server environments. Recent research has explored machine learning techniques, including regression models and ensemble methods, to improve forecasting accuracy. However, these models often require extensive feature engineering and fail to capture long-term dependencies inherent in complex server workloads. As a result, attention has shifted toward deep learning models that can automatically learn temporal representations directly from raw data.

2.2. Deep Learning Models for Anomaly Prediction

Deep learning has emerged as a powerful approach for anomaly prediction in time-series data due to its ability to model nonlinear and long-range temporal patterns. Recurrent neural networks, particularly Long Short-Term Memory (LSTM) networks, have been widely applied to detect anomalies in system metrics by predicting normal behavior and identifying deviations. More recently, Temporal Convolutional Networks (TCN) and Transformer-based architectures have demonstrated superior performance in modeling long sequences with improved training stability

and scalability. Autoencoders and sequence-to-sequence models have also been used to learn normal system behavior and flag anomalies based on reconstruction error. Despite these advances, challenges remain in applying deep learning models to real-world server telemetry, including interpretability, adaptability, and robustness to noise.

2.3. Existing Reliability and Telemetry Analytics Frameworks

Several commercial and open-source monitoring platforms provide telemetry analytics for server reliability management. Tools such as Prometheus, ELK Stack, and cloud-native monitoring services collect and visualize metrics, logs, and traces. While some platforms incorporate basic anomaly detection features, they often rely on heuristic or rule-based methods. Research frameworks have proposed integrating machine learning into monitoring pipelines, but many focus on offline analysis or specific failure scenarios. Moreover, existing frameworks often lack support for deep temporal modeling and explainability, limiting their adoption in production environments where trust and transparency are critical.

2.4. Gap Analysis

The existing literature highlights a gap between advanced deep learning research and practical server reliability management. While deep learning models show promise for time-series anomaly prediction, few studies provide a comprehensive framework that integrates data preprocessing, multivariate forecasting, anomaly detection, and explainability within an operational context. Additionally, comparative evaluations of different deep learning architectures on server telemetry data remain limited. This paper addresses these gaps by proposing a unified methodology that combines state-of-the-art time-series deep learning models with practical reliability objectives and deployment considerations.

3. PROBLEM FORMULATION

3.1. Server Telemetry Data Characteristics

Server telemetry data consists of continuous streams of time-stamped measurements collected from hardware, operating systems, and applications. These metrics are typically multivariate, high-frequency, and noisy, exhibiting seasonal patterns, sudden spikes, and long-term trends. Telemetry data may also contain missing values and irregular sampling intervals due to network delays or system failures. Understanding these characteristics is essential for designing effective prediction models, as they influence data preprocessing, feature extraction, and model selection. The multivariate nature of telemetry data requires models that can capture correlations across metrics and temporal dependencies over varying time scales.

3.2. Mathematical Formulation of Reliability Prediction

The reliability prediction problem can be formulated as a supervised or semi-supervised time-series learning task. Given a multivariate time-series representing historical server telemetry data over a fixed look-back window, the objective is to predict future metric values or a reliability-related signal, such as failure probability or anomaly score. Formally, the model learns a mapping from past

observations to future outcomes, minimizing a predefined loss function. Anomalies or reliability risks are identified when the predicted behavior deviates significantly from observed values or crosses learned thresholds. This formulation enables both forecasting-based anomaly detection and direct failure prediction.

3.3. Key Performance Indicators (KPIs) for Servers

Key performance indicators play a central role in server reliability management, as they reflect the operational health of systems. Common KPIs include CPU utilization, memory usage, disk I/O latency, network throughput, error rates, and response times. These indicators provide complementary views of system behavior and are often interdependent. For example, sustained high CPU usage may lead to increased response latency and error rates. Selecting appropriate KPIs and modeling their interactions is critical for accurately predicting reliability issues and minimizing false alarms.

3.4. Operational Objectives

The primary operational objective of proactive server reliability management is early fault detection, enabling timely intervention before failures impact users. Secondary objectives include reducing system downtime, minimizing false alerts, and optimizing resource utilization. Predictive models should provide sufficient lead time for corrective actions while maintaining high precision and recall. Additionally, the models should be robust, scalable, and interpretable to support real-world operational decision-making.

4. METHODOLOGY

4.1. Data Collection and Preprocessing

Data collection involves gathering telemetry data from multiple sources such as system logs, monitoring agents, and application-level metrics. These data sources provide complementary information about server behavior. Preprocessing is a critical step that includes handling missing values, removing noise, and aligning metrics across time. Normalization and scaling are applied to ensure that metrics with different ranges contribute equally to model training. Proper preprocessing improves model convergence and predictive accuracy while reducing the impact of outliers and measurement errors.

4.2. Feature Engineering

Feature engineering transforms raw telemetry data into meaningful representations that enhance model learning. Statistical features such as mean, variance, and rolling averages capture local trends and variability. Temporal features encode seasonality, trends, and time-based patterns, while derived metrics such as utilization ratios and rate of change provide additional insight into system dynamics. Although deep learning models can learn features automatically, carefully designed features can improve performance and interpretability, particularly in complex server environments.

4.3. Deep Learning Architectures

4.3.1. LSTM Based Model

LSTM networks are designed to capture long-term dependencies in sequential data through gated memory cells. In server telemetry analysis, LSTMs can model how historical system behavior influences future performance. Their ability to retain information over long time horizons makes them suitable for detecting gradual degradation and recurring patterns. However, LSTMs can be computationally expensive and challenging to scale for very long sequences.

4.3.2. Temporal Convolutional Networks (TCN)

Temporal Convolutional Networks use causal and dilated convolutions to model time-series data efficiently. TCNs offer parallel computation, stable gradients, and the ability to capture long-range dependencies with fewer parameters. In server reliability prediction, TCNs provide fast training and inference, making them suitable for real-time monitoring scenarios.

4.3.3. Transformer-Based Encoder Models

Transformer-based models leverage self-attention mechanisms to learn relationships between all time steps in a sequence. This allows them to capture complex temporal dependencies without recurrence or convolution. Transformers are particularly effective for multivariate telemetry data, as they can model interactions across metrics and time. However, their computational cost requires careful optimization for deployment.

4.3.4. Hybrid Model Variants

Hybrid models combine the strengths of different architectures, such as integrating convolutional layers with LSTM or attention mechanisms. These models aim to balance accuracy, efficiency, and interpretability. Hybrid approaches are especially useful in heterogeneous server environments with diverse workload patterns.

4.4. Training Strategy

Training strategies involve defining sliding windows and look-back periods to structure the time-series data for supervised learning. Appropriate loss functions such as mean absolute error, mean squared error, or Huber loss are selected based on robustness requirements. Regularization techniques and hyperparameter tuning are applied to prevent overfitting and improve generalization. Adaptive retraining strategies ensure that models remain effective as system behavior evolves.

4.5. Anomaly Thresholding and Event Detection

Anomaly detection is performed by comparing predicted and observed values to compute deviation scores. Thresholds are dynamically or statistically defined to distinguish normal fluctuations from abnormal behavior. Detected anomalies are mapped to reliability events, such as potential failures or performance degradation. Effective thresholding is essential to balance sensitivity and false alarm rates.

4.6. Explainable AI (XAI) Techniques for Model Interpretation

Explainability techniques are incorporated to provide insights into model predictions and build trust among operators. Methods such as attention visualization, feature importance analysis, and post-hoc explanation tools help identify which metrics and time periods contribute most to predictions. Explainable AI enhances transparency, facilitates debugging, and supports informed decision-making in operational environments.

5. EXPERIMENTAL SETUP

5.1. Dataset Description

The experimental evaluation of the proposed framework is conducted using server telemetry datasets that represent realistic operational environments. These datasets may consist of either real-world telemetry collected from production or staging servers, or carefully designed simulated telemetry streams that emulate common server behaviors and failure patterns. Real-world datasets typically include time-stamped metrics such as CPU utilization, memory consumption, disk I/O latency, network throughput, request rates, and error counts collected at regular intervals. Simulated datasets are often employed to supplement real data when failure events are rare, enabling controlled experimentation with known fault injection scenarios. Both types of datasets exhibit key time-series characteristics such as seasonality, workload variability, noise, and occasional missing values. The datasets are partitioned into training, validation, and testing subsets in a chronological manner to preserve temporal consistency and prevent data leakage.

5.2. Evaluation Metrics

To comprehensively assess model performance, multiple evaluation metrics are employed that capture both forecasting accuracy and operational effectiveness. Precision, recall, and F1-score are used to evaluate anomaly and failure detection performance, measuring the model's ability to correctly identify true reliability issues while minimizing false alarms. Root Mean Square Error (RMSE) is used to quantify forecasting accuracy by measuring the deviation between predicted and observed telemetry values. Mean Time to Detection (MTTD) is included as a critical operational metric that evaluates how early a model can detect failures relative to their actual occurrence. Together, these metrics provide a balanced evaluation of predictive accuracy, detection timeliness, and practical reliability impact.

5.3. Baseline Models

Baseline models are implemented to provide a reference point for evaluating the effectiveness of deep learning approaches. Traditional statistical models such as ARIMA are used to represent classical time-series forecasting techniques that rely on linear assumptions and stationarity. Machine learning baselines such as Random Forest and Support Vector Machines are included to capture nonlinear relationships using feature-based approaches. These models typically require handcrafted features and lack the ability to model long-term temporal dependencies inherently. Comparing deep learning models against these baselines highlights the advantages of end-to-end temporal

representation learning and demonstrates the added value of deep architectures in complex server telemetry scenarios.

5.4. Implementation Environment

The experimental framework is implemented using widely adopted open-source technologies to ensure reproducibility and scalability. Python serves as the primary programming language due to its extensive ecosystem for data analysis and machine learning. Deep learning models are implemented using frameworks such as TensorFlow or PyTorch, which provide efficient support for sequence modeling and GPU acceleration. Data preprocessing, feature extraction, and evaluation are conducted using standard scientific computing libraries. The experiments are executed on servers equipped with sufficient computational resources, and model training configurations are carefully documented to enable consistent and fair comparisons.

6. RESULTS AND ANALYSIS

6.1. Forecasting Accuracy Results

The forecasting accuracy results demonstrate the ability of the proposed deep learning models to predict future server behavior based on historical telemetry data. Models such as LSTM, TCN, and Transformer consistently achieve lower RMSE values compared to baseline approaches, indicating more accurate modeling of temporal dynamics. The results show that deep learning models effectively capture seasonality, workload shifts, and gradual performance degradation that traditional models fail to represent. Improved forecasting accuracy directly contributes to more reliable anomaly detection, as deviations from predicted behavior become more meaningful and actionable.

6.2. Comparative Performance across Architectures

A comparative analysis across different deep learning architectures reveals distinct strengths and trade-offs. LSTM models perform well in capturing long-term dependencies but may exhibit higher training time and sensitivity to hyperparameter settings. Temporal Convolutional Networks offer faster training and stable performance, particularly in high-frequency telemetry scenarios. Transformer-based models demonstrate superior performance in modeling multivariate dependencies and complex temporal interactions but require careful optimization to manage computational overhead. The comparative results provide insights into selecting appropriate architectures based on operational constraints and reliability objectives.

6.3. Early Failure Detection Performance

Early failure detection performance is evaluated by analyzing the models' ability to identify anomalies and reliability risks before actual failures occur. Deep learning models significantly reduce Mean Time to Detection compared to baseline methods, providing earlier warnings and increased lead time for intervention. High recall values indicate that the models successfully detect the majority of failure events, while precision metrics confirm a reduction in false positives. These results validate

the effectiveness of predictive modeling in transitioning from reactive monitoring to proactive reliability management.

6.4. Case Studies and Failure Scenarios

To illustrate real-world applicability, detailed case studies are presented that examine specific failure scenarios such as resource exhaustion, network congestion, and cascading service degradation. In each case, the model's predictions are analyzed alongside observed telemetry trends to demonstrate how early warning signals emerge prior to failure. These case studies highlight the interpretability and practical utility of the proposed framework, showing how operators can use predictions to make informed decisions and prevent service disruptions.

6.5. Ablation Studies

Ablation studies are conducted to assess the contribution of individual components within the proposed framework. By systematically removing or modifying elements such as feature engineering steps, model layers, or attention mechanisms, the impact of each component on overall performance is evaluated. The results show that multivariate input features, temporal context length, and architectural design choices significantly influence predictive accuracy and detection performance. Ablation analysis strengthens the validity of the proposed approach by demonstrating that performance gains are not attributable to a single factor.

6.6. Model Explainability Insights

Explainability analysis provides valuable insights into how deep learning models make predictions. Attention weights, feature importance measures, and gradient-based explanation techniques reveal which telemetry metrics and time intervals contribute most to predicted anomalies or failures. These insights help validate model behavior, align predictions with domain knowledge, and build trust among system operators. Explainable results also support root cause analysis by identifying potential sources of reliability issues, making the framework more actionable in operational environments.

7. DEPLOYMENT CONSIDERATIONS

7.1. Integration with Monitoring Systems

For practical adoption, the proposed framework is designed to integrate seamlessly with existing monitoring systems such as Prometheus and ELK Stack. Telemetry data collected by monitoring agents can be streamed directly into the prediction pipeline, enabling near real-time analysis. The integration ensures minimal disruption to existing workflows while enhancing monitoring capabilities with predictive intelligence. This compatibility facilitates deployment in diverse infrastructure environments, including on-premise, cloud, and hybrid systems.

7.2. Alerting and Automated Response

The predictive outputs of the model are integrated with alerting mechanisms to trigger notifications or automated responses when reliability risks are detected. Alerts are generated based

on anomaly scores or predicted failure probabilities, allowing operators to take corrective actions proactively. Automated responses such as resource scaling, workload redistribution, or service restarts can be initiated to mitigate potential failures. This closed-loop approach reduces manual intervention and enhances system resilience.

7.3. Scalability and Resource Efficiency

Scalability and efficiency are critical considerations for deploying predictive models in large-scale server environments. The framework is designed to handle high-volume telemetry streams while maintaining low inference latency. Model optimization techniques such as batching, model compression, and efficient architecture selection are employed to minimize computational overhead. These considerations ensure that the predictive system can scale with growing infrastructure without compromising performance.

7.4. Model Retraining and Online Learning

Server workloads and usage patterns evolve over time, necessitating continuous model adaptation. The framework supports periodic retraining using recent telemetry data to maintain predictive accuracy. Online learning strategies can be employed to update models incrementally as new data becomes available, reducing downtime and manual effort. Adaptive retraining ensures robustness against concept drift and sustains long-term reliability in dynamic operational environments.

8. DISCUSSION

8.1. Strengths of the Proposed Framework

The proposed time-series deep learning framework demonstrates several strengths that make it effective for proactive server reliability management. One of its primary advantages lies in its ability to model complex temporal dependencies and nonlinear relationships inherent in multivariate server telemetry data. By leveraging advanced deep learning architectures such as LSTM, Temporal Convolutional Networks, and Transformer-based models, the framework captures both short-term fluctuations and long-term trends that traditional monitoring approaches often overlook. The integration of forecasting-based anomaly detection enables early identification of reliability risks, providing valuable lead time for preventive actions. Furthermore, the inclusion of explainable AI techniques enhances transparency and trust, allowing system operators to understand model predictions and align them with operational knowledge. The modular design of the framework also supports extensibility, making it adaptable to diverse server environments and evolving workloads.

8.2. Limitations and Assumptions

Despite its strengths, the proposed framework is subject to several limitations and underlying assumptions. The effectiveness of deep learning models heavily depends on the quality and representativeness of historical telemetry data. In environments where failure events are rare or data is noisy and incomplete, model performance may be affected. The framework assumes that past system behavior contains sufficient information to predict future reliability issues, which may not

hold in cases of abrupt configuration changes or previously unseen failure modes. Additionally, deep learning models require significant computational resources for training and tuning, which may limit adoption in resource-constrained environments. While explainability techniques improve interpretability, they may not fully capture the internal complexity of deep models, leaving some predictions partially opaque.

8.3. Risks (False Positives and False Negatives)

A critical risk in predictive reliability systems is the occurrence of false positives and false negatives. False positives, where normal system behavior is incorrectly flagged as anomalous, can lead to alert fatigue, unnecessary interventions, and reduced trust in the monitoring system. Conversely, false negatives, where actual reliability issues go undetected, pose a more severe risk as they can result in unexpected failures and service downtime. The proposed framework mitigates these risks through adaptive thresholding, multivariate modeling, and continuous retraining; however, achieving an optimal balance between sensitivity and specificity remains a challenge. Careful calibration, domain-specific tuning, and human oversight are essential to manage these risks effectively in production environments.

8.4. Practical Implications for Site Reliability Engineering (SRE)

The proposed framework has significant practical implications for Site Reliability Engineering practices. By shifting from reactive alerting to predictive analytics, SRE teams can move toward a more proactive and preventive operational model. Early detection of reliability risks enables informed decision-making, such as preemptive scaling, traffic rerouting, or scheduled maintenance, thereby reducing downtime and improving service-level objectives. The framework also supports automation and continuous improvement, aligning with modern DevOps and SRE principles. Moreover, explainable predictions facilitate collaboration between engineers and automated systems, enhancing trust and operational efficiency.

9. CONCLUSION

9.1. Summary of Findings

This study demonstrates that time-series deep learning models are highly effective for proactive server reliability management. Through comprehensive experimentation and analysis, the results show that deep learning architectures significantly outperform traditional statistical and machine learning baselines in forecasting accuracy and early failure detection. The models successfully capture complex temporal patterns in server telemetry data, enabling timely identification of reliability risks. The integration of anomaly detection, explainability, and deployment considerations highlights the feasibility of applying these techniques in real-world operational environments.

9.2. Contributions to Software Reliability and Machine Learning Communities

The contributions of this work extend to both the software reliability and machine learning research communities. From a software engineering perspective, the paper provides a structured

framework for predictive reliability management that bridges the gap between monitoring and intelligent automation. From a machine learning standpoint, the study offers a comparative evaluation of multiple time-series deep learning architectures applied to real-world telemetry data, along with insights into explainability and operational deployment. The proposed methodology serves as a reference for future research at the intersection of AI and software reliability.

9.3. Future Work Directions

Future research can extend this work in several directions. Incorporating multimodal data sources such as logs, traces, and event streams alongside metric-based telemetry can provide a more holistic view of system behavior and improve prediction accuracy. Federated and privacy-preserving learning techniques offer promising avenues for collaborative reliability modeling across organizations without exposing sensitive data. Additionally, reinforcement learning can be explored to enable automated remediation strategies, where predictive models not only detect issues but also learn optimal actions to resolve them dynamically. These directions open new possibilities for fully autonomous and intelligent reliability management systems.

REFERENCES

- [1] J. Dean and L. A. Barroso, "The tail at scale," *Communications of the ACM*, vol. 56, no. 2, pp. 74–80, 2013.
- [2] C. Bishop, *Pattern Recognition and Machine Learning*, Springer, 2006.
- [3] S. Hochreiter and J. Schmidhuber, "Long short-term memory," *Neural Computation*, vol. 9, no. 8, pp. 1735–1780, 1997.
- [4] S. Bai, J. Z. Kolter, and V. Koltun, "An empirical evaluation of generic convolutional and recurrent networks for sequence modeling," *arXiv preprint arXiv:1803.01271*, 2018.
- [5] A. Vaswani et al., "Attention is all you need," *Advances in Neural Information Processing Systems*, pp. 5998–6008, 2017.
- [6] T. Li et al., "Deep learning for anomaly detection: A review," *ACM Computing Surveys*, vol. 54, no. 2, 2021.
- [7] M. Chen et al., "Machine learning for system reliability: A survey," *IEEE Transactions on Reliability*, vol. 69, no. 2, pp. 543–558, 2020.
- [8] R. J. Hyndman and G. Athanasopoulos, *Forecasting: Principles and Practice*, OTexts, 2018.
- [9] P. K. Agarwal et al., "Site reliability engineering: How Google runs production systems," O'Reilly Media, 2016.
- [10] D. He et al., "Experience report: System log analysis for anomaly detection," *IEEE International Conference on Software Reliability Engineering*, 2016.