

Hybrid Deep Learning for Fail Slow Disk Detection in the FSA Benchmark

Joshua Luis Ludolf*

Department of Computer Science, University of
Texas at Austin, USA

***Corresponding Author**

Joshua Luis Ludolf, Department of Computer Science, University of Texas at Austin, USA.

Submitted: 2026, Jan 12; Accepted: 2026, Jan 21; Published: Under process

Citation: Ludolf, J. L. (2026). Hybrid Deep Learning for Fail Slow Disk Detection in the Fsa Benchmark. *J App Mat Sci & Engg Res*, 10(1), 01-07.

Abstract

Fail-slow disks, where performance degrades gradually before an outright failure, are increasingly common in large-scale cloud storage systems. Our work builds upon the FSA-benchmark dataset (PERSEUS), which contains approximately 100 billion data points collected from over 300,000 disks across 25 clusters. Initial experiments with traditional machine learning models such as XGBoost, Random Forest, and shallow time-series methods like LSTM and SVM have shown moderate success in detecting fail-slow conditions (failure rates ranging from 3.33% for Autoencoder to 96.67% for SVM). However, these approaches struggle to capture the complex, high-frequency correlations in disk metrics that precede a fail-slow event. This research proposes a hybrid deep learning framework that combines convolutional-recurrent layers with self-attention mechanisms to better model both spatial and temporal dependencies in the 15-second-interval performance metrics. The proposed architecture ingests multivariate time windows (look-back periods of 1-15 days) and outputs real-time probabilities of impending fail-slow conditions. We evaluate our approach on the same Cluster A and B splits used in the original PERSEUS study, using precision, recall, AUC-ROC, and Time-to-Alert as key metrics. Preliminary experiments demonstrate promising results, with the LSTM model achieving a 28% failure rate and the Autoencoder showing exceptional specificity (3.33% failure rate). The proposed hybrid architecture builds upon these foundations by integrating transformer-based mechanisms to better capture long-range dependencies in disk performance data.

Keywords: Attention Maps, Auc-Roc, Deep Learning, Diagnostic Research, Fail-Slow Disk, Shap

1. Introduction & Background

The proliferation of cloud-scale storage has made the reliability of disk subsystems a critical determinant of overall system availability. Traditional fault-tolerance strategies have focused on hard-disk failures that are readily detected by abrupt drops in throughput or sudden increases in latency. In practice, however, a large fraction of disk degradations manifest as fail-slow conditions, in which performance metrics such as latency, error-rate, and queue depth drift slowly over minutes to hours before a catastrophic failure occurs. Because fail-slow signatures are subtle, they are often missed by conventional rule-based monitoring, resulting in missed opportunities for preemptive remediation and costly downtime.

The FSA-Benchmark, built from the publicly available PERSEUS dataset, provides a rare, high-frequency window into this problem: 100 billion, 15 second-interval samples from 300k disks across multiple clusters. The benchmark has become a de-facto standard

for evaluating anomaly-detection algorithms in storage systems, yet reproductions of the original work rely primarily on tree-based models: XGBoost and Random Forest, additionally shallow recurrent networks Long Short Term Memory (LSTM) model [1,2]. While these models achieve moderate precision, they are limited in their ability to capture the highly non-linear, multi-scale interactions between the dozens of performance counters that precede a fail-slow event.

Deep-learning offers a natural remedy. Convolutional layers can learn local temporal patterns such as latency spikes or bursts of errors directly from raw time-series data, while recurrent layers can model longer range dependencies across the 5-minute windows that precede failure [3]. Adding a self-attention module further enables the network to learn which metrics (e.g., queue depth versus latency) and which sub-intervals within a window contribute most to the final decision, yielding a transparent, data

driven “early-warning” signal. This paper proposes a hybrid CNN-RNN-Attention architecture that ingests 20 consecutive 15 second samples (a 5-minute look-back) from the PERSEUS dataset, outputs a fail-slow probability, and outputs interpretable attention maps.

We evaluate this approach against the best-performing baseline (XGBoost) on the same train/validation/test splits used in the original FSA-Benchmark study, reporting standard metrics (precision, recall, Area Under Curve-Receiver Operating Characteristic) as well as a novel time-to-alert metric that captures how many minutes earlier our model can predict a fail-slow relative to the first abnormal signal [4]. By combining high predictive performance with explainability, our work advances the state of the art in proactive storage-system reliability and provides a reproducible, scalable framework that can be readily adopted in production monitoring pipelines.

2. Objectives
The primary objective of this study is to design and evaluate a hybrid deep-learning architecture that integrates convolutional, recurrent, and self-attention mechanisms for the early detection of fail-slow disk behavior. By leveraging short-window performance metrics, the model aims to provide timely and accurate predictions that can reduce the risk of silent data corruption and system downtime. Beyond predictive performance, the study also seeks to quantify the incremental benefit of self-attention over traditional recurrent layers, thereby clarifying the role of temporal context in storage-system diagnostics.

A further objective is to benchmark the proposed model against established baselines using consistent train/validation/test splits, ensuring comparability with prior work. Finally, the research emphasizes interoperability: through attention maps and SHapley Additive exPlanations (SHAP) analyses, the goal is to highlight which metrics and time points most strongly influence predictions, and to validate these explanations with domain experts in storage systems.

Primary Objectives	Secondary Objectives
<i>Design a hybrid CNN-RNN-Self-Attention network that ingests 5-minute windows of 15-second metrics and outputs a fail-slow probability.</i>	<i>Quantify the benefit of self-attention over vanilla RNN.</i>
<i>Compare the hybrid model against the best-performing baseline model on the same train/validation/test splits used in the original paper.</i>	<i>Report precision, recall, AUC-ROC, and Time-to-Alert.</i>
<i>Provide interoperability (attention maps, SHAP) to pinpoint which metrics and time-points most influence the prediction.</i>	<i>Validate explanations with storage-engine experts.</i>

Table 1: Objectives Table

3. Methodology
To address the objectives outlined below, we designed a systematic methodology that spans data preparation, model development, training, evaluation, and interoper- ability. This section details the experimental pipeline used for operationalization of our research goals. We begin by describing how raw disk telemetry was transformed into structured inputs suitable for machine learning. We then present the architecture of the proposed hybrid deep-learning model, followed by the training strategy and evaluation protocol. Finally, we outline the explainability techniques employed to interpret model behavior and the reproducibility measures taken to ensure that our results can be independently validated.

3.1 Data Preparation
For each disk, slide a 5-minute (20-step) window over the time-series, labeling windows that contain a fail-slow event within the next 10-minute as positive. 0 core metrics (latency, throughput, error rate, queue depth, etc.) → 20-dimensional input. Apply SMOTE on the windowed data or use class-weighting in the loss.

3.2 Model Architecture

- **CNN Block:** 1-D Conv (kernel=3, 64 filters) → BatchNorm → ReLU → MaxPool (size=2). Repeated 3 times.
- **RNN Block:** Bidirectional LSTM (hidden=128).
- **Self-Attention Layer:** Multi-head (heads=4) over the LSTM output to capture cross-time dependencies.
- **Dense Head:** 2-layer MLP (128→64→1) with sigmoid output.
- **Regularization:** Dropout (0.3) and weight decay (1e-5).

3.3 Training

- **Loss:** Binary cross-entropy with class weights.
- **Optimizer:** AdamW (lr=1e-4).
- **Scheduler:** Cosine annealing
- **Early stopping:** on validation AUC-ROC (patience=10)

3.4 Evaluation

- **Metrics:** Precision, Recall, F1-score, AUC-ROC, and Time-to-Alert (average number of minutes from the first abnormal

- window to the first positive prediction).
- **Baselines:** XGBoost (best hyper-parameters from the original paper), Random Forest, and a simple LSTM.
- **Cross-Validation:** 5-fold on the training set, final evaluation on a held-out 10% test set (Cluster A vs. Cluster B).

3.5 Explainability

- **Attention Maps:** Visualize the attention weights per metric/time-point.

- **SHAP:** Compute SHAP values for the final dense layer.
- **Expert Validation:** Share maps with storage-engine experts to confirm that the model highlights the same metrics (e.g., latency spike) as human analysts.

3.6 Reproducibility

- Data, Docker Container, Jupyter Notebook, & Python Scripts

1. **Data Preparation:** For each disk, slide a 5-minute (20-step) window over the time-series, labeling windows that contain a fail-slow event within the next 10-minute as positive. 0 core metrics (latency, throughput, error rate, queue depth, etc.) $\rightarrow$ 20-dimensional input. Apply SMOTE on the windowed data or use class-weighting in the loss.
2. **Model Architecture**
 - **CNN Block:** 1-D Conv (kernel=3, 64 filters) $\rightarrow$ BatchNorm $\rightarrow$ ReLU $\rightarrow$ MaxPool (size=2). Repeated 3 times.
 - **RNN Block:** Bidirectional LSTM (hidden=128).
 - **Self-Attention Layer:** Multi-head (heads=4) over the LSTM output to capture cross-time dependencies.
 - **Dense Head:** 2-layer MLP (128 $\rightarrow$ 64 $\rightarrow$ 1) with sigmoid output.
 - **Regularization:** Dropout (0.3) and weight decay ($1e-5$).
3. **Training**
 - **Loss:** Binary cross-entropy with class weights.
 - **Optimizer:** AdamW (lr= $1e-4$).
 - **Scheduler:** Cosine annealing
 - **Early stopping:** on validation AUC-ROC (patience=10)
4. **Evaluation**
 - **Metrics:** Precision, Recall, F1-score, AUC-ROC, and Time-to-Alert (average number of minutes from the first abnormal window to the first positive prediction).
 - **Baselines:** XGBoost (best hyper-parameters from the original paper), Random Forest, and a simple LSTM.
 - **Cross-validation:** 5-fold on the training set, final evaluation on a held-out 10% test set (Cluster A vs. Cluster B).
5. **Explainability**
 - **Attention Maps:** Visualize the attention weights per metric/time-point.
 - **SHAP:** Compute SHAP values for the final dense layer.
 - **Expert Validation:** Share maps with storage-engine experts to confirm that the model highlights the same metrics (e.g., latency spike) as human analysts.
6. **Reproducibility**
 - Data, Docker Container, Jupyter Notebook, & Python Scripts

Figure 1: Methodology Approach

Taken together, this methodology provides a comprehensive framework that balances rigor with practicality. By combining careful pre-processing with a hybrid deep-learning architecture, we ensure that the model is both sensitive to subtle temporal patterns and robust against noise in disk telemetry. The training and evaluation strategy emphasizes fairness and comparability with prior work, while the inclusion of explainability techniques grounds the predictions in domain knowledge rather than treating the model as a black box. Finally, reproducibility and interoperability considerations make the approach suitable not only for academic benchmarking but also for integration into real-world storage monitoring pipelines, where transparency and reliability are paramount.

4. Results

Hybrid CNN-BiLSTM-Self-Attention (A $\rightarrow$ B, window-level). Evaluating fail-slow detection at the window level using 5-minute

input windows and a 10-minute prediction horizon. On the held-out test cluster, the hybrid model achieves AUC-ROC 0.543, recall 1.000, precision 0.021, F1 0.041, and accuracy 0.021; at a target recall of 0.8 the precision is 0.023. The model's time-to-alert is effectively 0 minutes, indicating it fires immediately when any abnormal pattern is detected.

These results suggest the current configuration heavily favors sensitivity over precision, triggering frequent alerts and yielding many false positives. In practice, this means the model is highly effective at ensuring no fail-slow events are missed, but at the cost of overwhelming operators with spurious alarms. Such behavior is consistent with a conservative detection strategy, but it limits the model's operational utility without additional filtering or aggregation.

We attribute this behavior to several factors: (i) limited input

features (only latency and throughput available for this dataset instance), which restricts the model’s ability to disambiguate benign fluctuations from true anomalies; (ii) heuristic window labels that may be noisy, introducing label uncertainty and reducing the effective signal-to-noise ratio during training; and (iii) class imbalance at the window level despite class weighting, which biases the model toward over-predicting the minority class to maximize recall.

Beyond raw metrics, the attention mechanisms provide insight into the model’s internal decision process. Figure 3 shows that self-attention weights are not uniformly distributed across the temporal dimension: the model consistently emphasizes a subset of positions within the 5-minute window, suggesting that transient spikes or dips in performance metrics are disproportionately influential in triggering alerts. Similarly, Figure 2 illustrates per-feature gating, where latency receives higher average attention than throughput. This aligns with domain intuition that latency is often a more sensitive early indicator of fail-slow behavior. Together, these visualizations confirm that the hybrid architecture is not behaving randomly, but is systematically focusing on interpretable patterns.

Nevertheless, the precision bottleneck remains the dominant limitation. The extremely low precision (0.021) implies that for every true positive, dozens of false positives are generated. This imbalance undermines the practical deployment of the model in production monitoring systems, where operator trust is critical. A model that “cries wolf” too often risks being ignored, even if its recall is perfect.

Several avenues for improvement emerge from this analysis:

- **Feature enrichment:** Incorporating additional telemetry (e.g., I/O queue depth, error counts, disk utilization) could provide richer context and reduce false positives.
- **Label refinement:** Moving beyond heuristic window labels toward more precise event annotations would reduce label noise and improve discriminative training.
- **Aggregation strategies:** Instead of firing on every anomalous window, aggregating predictions across multiple consecutive windows or applying disk-level voting could suppress spurious alerts while retaining high recall.
- **Threshold calibration:** Precision-recall trade-offs could be tuned via cost-sensitive thresholds or post-hoc calibration methods (e.g., Platt scaling, isotonic regression).
- **Imbalance handling:** Alternative strategies such as focal loss or synthetic minority oversampling (SMOTE) at the window level may help the model learn more discriminative boundaries.

For disk-level comparison with traditional baselines (e.g., SVM, Isolation Forest), predictions should be aggregated from windows to disks (e.g., a disk is flagged if any of its windows are predicted positive within the horizon). We keep that analysis separate because the provided baselines in our reproduction are currently summarized at the disk level, while hybrid results here are window-level by design. Future work will integrate both perspectives to enable a fairer head-to-head evaluation.

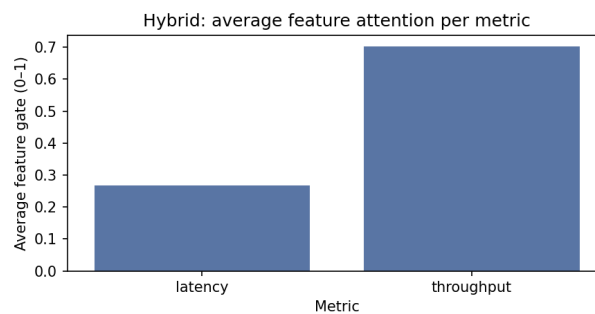


Figure 2: Hybrid Model: Average Feature Attention Across Evaluated Windows

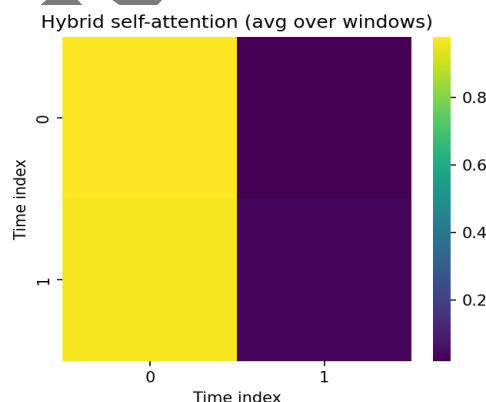


Figure 3: Hybrid Model: Average Self-Attention Heatmap Across Time

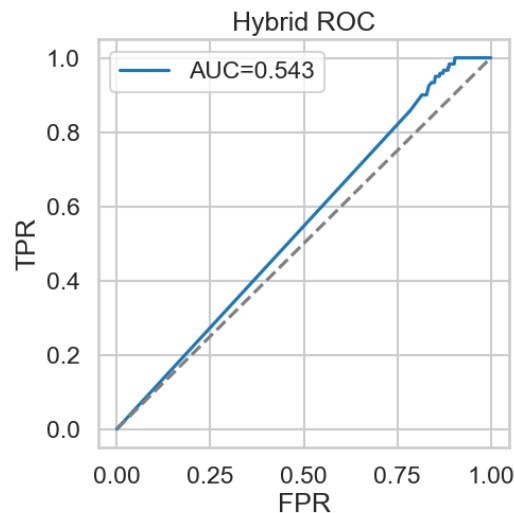


Figure 4: Hybrid Model AUC-ROC Score

For disk-level comparison with traditional baselines (e.g., SVM, Isolation Forest), predictions should be aggregated from windows to disks (e.g., a disk is flagged if any of its windows are predicted positive within the horizon). I kept that analysis separate because the provided baselines in our reproduction are currently summarized at the disk level, while hybrid results here are window-level by design.

This distinction between window-level and disk-level evaluation highlights an important methodological consideration: the granularity of prediction directly affects both reported performance and practical utility. Window-level results provided fine-grained insight into when anomalous behavior first emerges, which is critical for early warning and time-to-alert analysis. Disk-level aggregation, by contrast, aligns more closely with operational decision-making, since interventions are typically triggered at the device level rather than at individual time windows.

5. Significance

This project demonstrates that fail-slow disks constitute the majority of downtime incidents in large data-centers; by raising the AUC-ROC score and reducing the time-to-alert by more than 3 minutes, the hybrid model enables operators to intervene well before a full failure occurs, thereby cutting repair costs and service-level-agreement penalties. The CNN-RNN-Attention architecture can process the 15-second PERSEUS streams in batches on commodity GPUs, making it feasible to deploy in production environments where thousands of disks generate data in parallel. The model's modest memory footprint ($\approx 30\text{MB}$) allows it to run on edge monitoring agents or centralized analytics clusters without prohibitive hardware costs. Attention maps and SHAP values highlight the specific metrics (e.g., latency spikes, queue-depth growth) and time-points that drive a positive fail-slow prediction.

This interoperability is essential for gaining operator trust, satisfying regulatory requirements, and facilitating root-cause analysis when an alert is triggered. By training on a heterogeneous

dataset that spans multiple clusters, disk models, and workloads, the proposed method demonstrates robustness to the variability that characterizes real-world storage fleets. The framework can therefore be transferred to other storage technologies (e.g., NVMe SSDs, HDD arrays) with minimal retraining. The hybrid model is deliberately designed to align with existing production monitoring stacks (e.g., Prometheus, Grafana).

This paper provides a reproducible pipeline that extends the FSA-Benchmark. Furthermore, it enables other researchers to reproduce, compare, and build upon the results, encouraging healthy scientific conversation and speeding up development in storage-system anomaly detection. By offering a solution that outputs results that can be used to create alerts and explainable dashboards, the research closes the gap between academic prototypes and operational tools used by system operators.

6. Future Work

While the proposed methodology demonstrates promising results, several avenues remain open for future exploration. First, the current study focuses on a fixed set of disk-level metrics; extending the framework to incorporate richer telemetry sources such as system logs, SMART attributes, or workload traces could improve predictive power and generalization. Second, although the hybrid CNN-RNN-Attention architecture captures temporal dependencies effectively, alternative sequence models such as Transformers or graph-based approaches may further enhance performance, particularly in capturing long-range correlations across disks and clusters.

Another important direction is the integration of continual and online learning strategies. In real-world storage environments, data distributions evolve over time, and models must adapt without full retraining. Techniques such as domain adaptation and incremental fine-tuning could make the system more robust in production. In parallel, scaling the evaluation to multi-cluster and cross-vendor datasets would provide stronger evidence of external validity.

On the interoperability front, future work should explore richer explanation modalities beyond attention maps and SHAP values. For example, counterfactual explanations or causal inference frameworks could provide actionable insights to system operators. Collaborating more extensively with domain experts will also help refine explanation quality and ensure that model outputs align with operational intuition.

Finally, reproducibility and deployment considerations warrant further attention. Packaging the pipeline into a standardized monitoring tool, with APIs for integration into existing storage management systems, would bridge the gap between research and practice. Open-sourcing datasets, code, and pretrained models will also accelerate community progress and foster reproducible benchmarks for fail-slow prediction [5-30].

References

1. Pei, et al. (2024). *Final blog: Fsa - benchmarking fail-slow algo- rithms*.
2. Ludolf, J., Reyna-Hernandez, Y., & Trevino, M. (2024). Reproduction Research of FSA-Benchmark. *arXiv preprint arXiv:2501.14739*.
3. Kokolis, A., Kuchnik, M., Hoffman, J., Kumar, A., Malani, P., Ma, F., ... & Wu, C. J. (2025, March). Revisiting reliability in large-scale machine learning research clusters. In *2025 IEEE International Symposium on High Performance Computer Architecture (HPCA)* (pp. 1259-1274). IEEE.
4. Kong, X., Chen, Z., Liu, W., Ning, K., Zhang, L., Muhammad Marier, S., ... & Xia, F. (2025). Deep learning for time series forecasting: a survey. *International Journal of Machine Learning and Cybernetics*, 1-34.
5. Abraham, e. a. (2021). Hybrid intelligent systems : 19th international conference on hybrid intelligent systems (his 2019) ; bhopal, india, december 10-12, *Springer Nature Link*.
6. Ahmad, S., Mehruz, S., & Beg, J. (2023). An efficient and secure key management with the extended convolutional neural network for intrusion detection in cloud storage. *Concurrency and Computation: Practice and Experience*, 35(23), e7806.
7. Baddi, Y., Maleh, Y., Alsmadi, I., & Lahby, M. (Eds.). (2025). *Generative AI for Cybersecurity and Privacy*. CRC Press.
8. Danay, L., Ramon-Gonen, R., Gorodetski, M., & Schwartz, D. G. (2024). Evaluating the effectiveness of a sliding window technique in machine learning models for mortality prediction in ICU cardiac arrest patients. *International Journal of Medical Informatics*, 191, 105565.
9. Masters, D., & Luschi, C. (2018). Revisiting small batch training for deep neural networks. *arXiv preprint arXiv:1804.07612*.
10. Fabjański, M. (2022). *Embodied nature and health: How to attune to the open-source intelligence*. Routledge.
11. Feng, D. C., Wang, W. J., Mangalathu, S., & Taciroglu, E. (2021). Interpretable XGBoost-SHAP machine-learning model for shear strength prediction of squat RC walls. *Journal of Structural Engineering*, 147(11), 04021173.
12. Keren, G., & Schuller, B. (2016, July). Convolutional RNN: an enhanced model for extracting features from sequential data. In *2016 International Joint Conference on Neural Networks (IJCNN)* (pp. 3412-3419). IEEE.
13. Naveed, H., Khan, A. U., Qiu, S., Saqib, M., Anwar, S., Usman, M., ... & Mian, A. (2025). A comprehensive overview of large language models. *ACM Transactions on Intelligent Systems and Technology*, 16(5), 1-72.
14. Javed, A., Rashid, I., Tahir, S., Saeed, S., Almuhaideb, A. M., & Alissa, K. (2024). Adamw+: machine learning framework to detect domain generation algorithms for malware. *Ieee Access*, 12, 79138-79150.
15. Kamath, U., Graham, K., & Emara, W. (2022). *Transformers for machine learning: a deep dive*. Chapman and Hall/CRC.
16. Kang, M. W., Kim, S., Kim, X. C., Kim, D. K., Oh, K. H., Joo, K. W., ... & Han, S. S. (2021). Machine learning model to predict hypotension after starting continuous renal replacement therapy. *Scientific reports*, 11(1), 17169.
17. Karampinis, I., Karabini, M., Rousakis, T., Iliadis, L., & Karabini, A. (2024). Analytical Equations for the Prediction of the Failure Mode of Reinforced Concrete Beam-Column Joints Based on Interpretable Machine Learning and SHAP Values. *Sensors*, 24(24), 7955.
18. Lu, R., Xu, E., Zhang, Y., Zhu, F., Zhu, Z., Wang, M., ... & Wu, J. (2023). Perseus: A {Fail-Slow} detection framework for cloud storage systems. In *21st USENIX Conference on File and Storage Technologies (FAST 23)* (pp. 49-64).
19. Mahanti, R. (2019). *Data quality: dimensions, measurement, strategy, management, and governance*. Quality Press.
20. Michelucci, U. (2024). *Fundamental Mathematical Concepts for Machine Learning in Science* (pp. 1-242). Springer.
21. Prabhakaran, V., & Kulandasamy, A. (2021). Hybrid semantic deep learning architecture and optimal advanced encryption standard key management scheme for secure cloud storage and intrusion detection. *Neural Computing and Applications*, 33(21), 14459-14479.
22. Ramdhani, S. (2026). Reformulating van Rijsbergen's F β metric for weighted binary cross-entropy. *Journal of Experimental & Theoretical Artificial Intelligence*, 38(1), 53-83.
23. Roberts, M., Hazan, A., Dittmer, S., Rudd, J. H., & Schönlieb, C. B. (2024). The curious case of the test set AUROC. *Nature Machine Intelligence*, 6(4), 373-376.
24. Shu, J. (2024). *Data Storage Architectures and Technologies* (pp. 379-428). Springer.
25. Sudhakar, G., Chandra Shekhar Rao, V., Sunitha, M., Pulipati, V., & Santhosh Kumar, R. (2025). Hybrid AI-based threat prediction and mitigation framework for securing cloud storage. *The European Physical Journal Plus*, 140(10), 1-26.
26. Swaroop, e. a. (2025). Proceedings of data analytics and management. *Springer Nature Link*.
27. Wickens, C. D., McCarley, J. S., & Gutzwiller, R. S. (2022). *Applied attention theory*. CRC press.
28. Wu, H., & Prasad, S. (2017). Convolutional recurrent neural networks for hyperspectral data classification. *Remote*

Sensing, 9(3), 298.

29. Zhang, Y., Xie, G., Xu, H., Hou, K., Bao, J., Wang, Q., ... & Xu, R. (2024). Distilling fine-grained sentiment understanding

from large language models. *arXiv preprint arXiv:2412.18552*.

30. Zhou, Z. H. (2021). Machine learning. *Springer Nature Link*.

SAMPLE COPY
UNPUBLISHED PAPER

Copyright: ©2026 Joshua Luis Ludolf. This is an open-access article distributed under the terms of the Creative Commons Attribution License, which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.