

Finops & Resource Efficiency: Predictive Autoscaling Using Time-Series Analysis to Reduce Cloud Waste in Eks Clusters

Pavan Madduri

1221 FARMER CIR, SOUTH ELGIN , ILLINIOS 60177, USA.

pavanmadduri27@gmail.com

ABSTRACT

Cloud infrastructure costs represent significant operational expenditure for modern enterprises, with Kubernetes clusters on Amazon EKS often exhibiting substantial resource waste due to reactive autoscaling approaches that provision capacity only after demand materializes. This research develops a predictive autoscaling framework leveraging time-series analysis to proactively scale EKS cluster resources based on forecasted demand, reducing cloud waste while maintaining application performance. Through systematic analysis of workload patterns across twelve production EKS deployments spanning e-commerce, financial services, and media streaming organizations, we identify recurring temporal patterns enabling accurate demand forecasting. Our framework employs ARIMA and Prophet models for time-series prediction combined with custom Kubernetes controllers that translate forecasts into proactive scaling actions. Implementation across production environments processing 2.8 million requests daily demonstrates 34% reduction in infrastructure costs through elimination of over-provisioning, 41% decrease in autoscaling-induced performance degradation from pre-scaling before demand spikes, and 23% improvement in resource utilization efficiency. The research contributes both theoretical understanding of cloud workload predictability and practical frameworks enabling FinOps teams to optimize Kubernetes spending without compromising reliability. Our findings reveal that predictive autoscaling proves most effective for workloads with strong temporal patterns like business hours activity, scheduled batch processing, and recurring traffic events, while providing limited benefits for purely stochastic workloads.

KEYWORDS: FinOps, Cloud Cost Optimization, Kubernetes Autoscaling, Time-Series Forecasting, EKS, Resource Efficiency, Predictive Scaling

INTRODUCTION

Cloud computing has transformed IT infrastructure from capital expenditure to operational expenditure, enabling organizations to provision resources on-demand without upfront hardware investments. However, this flexibility introduces new challenges around cost management and resource efficiency. Industry reports indicate that 30-40% of cloud spending goes to waste through over-provisioned resources, idle capacity, and inefficient autoscaling strategies (Kumar and Roberts, 2023). For organizations running containerized workloads on Kubernetes platforms like Amazon Elastic Kubernetes Service, this waste manifests through clusters sized for peak capacity that remain underutilized during normal operations.

Traditional Kubernetes autoscaling employs reactive approaches where the Horizontal Pod Autoscaler monitors current CPU or memory utilization and adjusts replica counts when thresholds are exceeded. Similarly, Cluster Autoscaler provisions additional nodes when pods cannot be scheduled due to resource constraints. While these mechanisms prevent capacity shortfalls, they introduce several inefficiencies. First, scaling actions occur after demand increases, creating temporary performance degradation as new resources provision and initialize. Second, reactive scaling often over-provisions resources as safety buffers since organizations prefer excess capacity over potential outages. Third, scale-down operations lag behind demand decreases, leaving expensive resources idle (Thompson and Chen, 2024).

Financial Operations (FinOps) has emerged as a discipline combining financial accountability, cloud cost management, and technical optimization to maximize cloud value. FinOps principles emphasize right-sizing resources, eliminating waste, and aligning spending with business outcomes. Yet most FinOps initiatives rely

on manual analysis and periodic optimization rather than continuous automated optimization integrated into infrastructure operations (Anderson et al., 2023).

The core insight driving this research is that many cloud workloads exhibit predictable temporal patterns enabling proactive resource planning. E-commerce platforms experience daily traffic cycles following customer behavior patterns, with peaks during lunch hours and evenings while remaining quiet overnight. Financial services workloads spike at market open and close with consistent daily patterns. Media streaming exhibits evening peaks as users consume content after work. Batch processing jobs run on scheduled intervals creating perfectly predictable demand. These temporal patterns can be captured through time-series analysis and used to forecast future resource needs with sufficient accuracy for proactive scaling (Williams and Martinez, 2024).

This research develops a comprehensive predictive autoscaling framework specifically designed for EKS clusters that leverages time-series forecasting to anticipate demand and proactively scale resources before usage increases materialize. The framework addresses fundamental questions: What time-series models provide sufficient forecast accuracy for production autoscaling? How should forecasts translate into scaling decisions considering forecast uncertainty? What workload characteristics indicate predictability suitable for predictive autoscaling? How can predictive and reactive autoscaling integrate to handle both expected patterns and unexpected spikes?

Our contributions extend beyond algorithmic improvements to encompass complete operational frameworks including monitoring infrastructure capturing historical patterns, model training and validation pipelines, forecast-to-scaling-decision logic, and integration with existing Kubernetes autoscaling mechanisms. We validate the framework through production deployments demonstrating measurable cost reductions and performance improvements.

OBJECTIVES

- **Primary Objective:** Develop and validate a predictive autoscaling framework for EKS clusters that leverages time-series analysis to reduce cloud infrastructure costs while maintaining application performance and reliability.
- **Secondary Objective 1:** Identify time-series forecasting models that provide sufficient accuracy for production autoscaling decisions across diverse workload patterns including web applications, batch processing, and API services.
- **Secondary Objective 2:** Design scaling decision logic that translates demand forecasts into proactive resource provisioning while accounting for forecast uncertainty and maintaining safety margins.
- **Secondary Objective 3:** Evaluate framework effectiveness through production deployments, measuring cost reduction, resource utilization improvement, and performance impact compared to reactive autoscaling baselines.

SCOPE OF STUDY

The research encompasses:

- **Platform Scope:** Framework targets Amazon EKS specifically, though architectural principles apply to other managed Kubernetes platforms including AKS and GKE with platform-specific adaptations.
- **Workload Scope:** Analysis includes web applications, REST APIs, batch processing jobs, and data processing pipelines rather than focusing on single application types.
- **Forecasting Scope:** Time-series models evaluated include ARIMA, Prophet, LSTM networks, and ensemble approaches while excluding non-temporal prediction methods like regression on static features.
- **Metric Scope:** Evaluation measures infrastructure costs, resource utilization efficiency, application performance (latency, error rates), and autoscaling responsiveness.

- **Exclusions:** The study does not address application-level optimization, database scaling, or storage optimization which involve distinct technical approaches beyond compute autoscaling.

LITERATURE REVIEW

4.1 Cloud Cost Optimization and FinOps

FinOps emerged as organizations recognized that cloud's pay-per-use model requires continuous cost management rather than periodic optimization. The FinOps Foundation defines core principles including teams taking ownership of cloud usage, decisions driven by business value of cloud, and everyone taking responsibility for cloud usage. Yet translating these principles into technical practices remains challenging (Kumar and Roberts, 2023).

Cloud cost optimization research has examined various dimensions including right-sizing where resources match actual requirements, reserved capacity purchasing trading flexibility for discounts, and spot instance utilization accepting interruption risk for lower costs. However, most approaches treat optimization as periodic activity rather than continuous automated process. Autoscaling represents critical cost management mechanism but remains predominantly reactive (Morrison and Lee, 2024).

4.2 Kubernetes Autoscaling Mechanisms

Kubernetes provides multiple autoscaling mechanisms addressing different resource types. Horizontal Pod Autoscaler scales pod replica counts based on observed CPU, memory, or custom metrics. Vertical Pod Autoscaler adjusts resource requests and limits for individual pods. Cluster Autoscaler provisions or deprovisions worker nodes based on pod scheduling needs. These mechanisms operate independently through different controllers, creating coordination challenges (Harrison and Taylor, 2023).

Traditional autoscaling employs reactive threshold-based policies. When current utilization exceeds target thresholds, scaling actions trigger to add capacity. When utilization drops below lower thresholds for sustained periods, scale-down occurs. This reactive approach introduces inherent lag between demand changes and capacity adjustments. For rapidly changing workloads, this lag causes performance degradation during scale-up and resource waste during delayed scale-down (Thompson and Chen, 2024).

4.3 Time-Series Forecasting for Cloud Workloads

Time-series analysis provides statistical methods for modeling temporal data and forecasting future values. Classical approaches like ARIMA model auto-regressive and moving average components capturing trend and seasonality. Modern techniques including Facebook's Prophet handle multiple seasonality patterns and holiday effects. Deep learning approaches using LSTM networks learn complex temporal dependencies though requiring substantial training data (Patel and Wilson, 2024).

Research applying time-series forecasting to cloud workload prediction has shown promising results. Studies demonstrate that cloud workloads often exhibit daily, weekly, and monthly patterns enabling forecast accuracy sufficient for capacity planning. However, most research remains theoretical or uses synthetic workloads rather than validating predictions through actual autoscaling in production environments (Chen and Kumar, 2023).

Forecast accuracy requirements for autoscaling differ from pure prediction tasks. Perfect forecasts are unnecessary if prediction errors fall within safety margins. Conversely, even accurate average predictions may be insufficient if forecast confidence intervals are too wide for reliable scaling decisions. The relationship between forecast quality and autoscaling effectiveness requires empirical investigation (Gupta et al., 2024).

4.4 Predictive Autoscaling Research

Academic research on predictive autoscaling has explored various approaches. Rule-based systems use simple heuristics like "if traffic increased at 9 AM yesterday, preemptively scale at 8:55 AM today." Machine learning

approaches train models on historical metrics to predict future resource needs. Reinforcement learning frames autoscaling as sequential decision-making problem where policies learn optimal actions through trial and error (Roberts and Zhang, 2024).

However, substantial gaps remain between research and production deployment. Most published work evaluates approaches through simulation or controlled experiments rather than production systems. Research often focuses narrowly on prediction accuracy without addressing operational concerns like model retraining, forecast freshness, integration with existing systems, and handling prediction failures (Jackson and Lee, 2023).

4.5 Research Gaps

Literature reveals several critical gaps this research addresses. First, comprehensive frameworks integrating time-series forecasting with Kubernetes autoscaling in production EKS environments are absent. Existing work examines components in isolation without end-to-end implementation.

Second, empirical validation of predictive autoscaling in diverse production workloads is limited. Most research uses synthetic traces or narrow workload types, leaving uncertainty about real-world effectiveness across different application patterns.

Third, cost-performance trade-off analysis comparing predictive and reactive autoscaling under realistic conditions needs development. Organizations require evidence quantifying cost savings, performance impacts, and operational overhead before adopting predictive approaches.

This research fills these gaps through comprehensive framework development validated across production EKS deployments with diverse workload characteristics.

RESEARCH METHODOLOGY

This research employed design science methodology developing and validating predictive autoscaling artifacts through iterative cycles of design, implementation, and evaluation.

Workload Analysis Phase: Historical metrics from twelve production EKS deployments were collected over three-month periods capturing CPU utilization, memory consumption, request rates, and pod counts at one-minute granularity. Systematic analysis identified temporal patterns including daily cycles, weekly patterns, monthly trends, and special event impacts.

Model Development: Multiple time-series forecasting approaches were implemented including ARIMA for stationary series with trend and seasonality, Prophet for series with multiple seasonal patterns and holiday effects, and LSTM networks for complex non-linear patterns. Models were trained on historical data with hyperparameter optimization through grid search and cross-validation.

Framework Implementation: A Kubernetes custom controller was developed that periodically generates demand forecasts, translates predictions into scaling decisions, and issues scaling commands to Horizontal Pod Autoscaler and Cluster Autoscaler. The controller integrates with Prometheus for metrics collection and maintains forecast models through automated retraining pipelines.

Production Validation: Framework deployed across six production EKS clusters representing diverse workload types. Each deployment ran for eight-week evaluation periods with A/B testing comparing predictive autoscaling against reactive baselines. Comprehensive monitoring captured costs, performance metrics, and operational incidents.

Comparative Analysis: Statistical analysis compared predictive and reactive approaches across infrastructure costs calculated from EC2 instance hours and associated charges, resource utilization measured as average

CPU and memory usage percentages, performance metrics including request latency and error rates, and scaling efficiency measured as time to achieve target capacity.

PREDICTIVE AUTOSCALING FRAMEWORK

6.1 Temporal Pattern Identification

Analysis of production workloads revealed distinct temporal pattern classes. Daily cycles exhibited consistent intraday variation with traffic peaks during business hours and troughs overnight. Weekly patterns showed different activity levels across weekdays versus weekends. Monthly trends reflected billing cycles, payroll processing schedules, or seasonal business variations. Event-driven patterns corresponded to marketing campaigns, product launches, or scheduled batch jobs.



Figure 1: Workload Temporal Patterns

This multi-panel time-series visualization displays four representative workload patterns over one week. The top panel shows an e-commerce API exhibiting strong daily cycles with request rates (y-axis, thousands of requests per minute) ranging from 2K RPM at 3 AM to 18K RPM at peak periods around 1 PM and 7 PM each day. The pattern repeats consistently across weekdays (Monday-Friday) with slight variations, while weekends show different patterns with later peaks and lower overall volume. The second panel displays financial services batch processing with weekly patterns characterized by zero activity on weekends and concentrated processing windows at market open (9:30 AM) and close (4:00 PM) on weekdays, creating distinct spikes reaching 25K RPM during processing windows and near-zero activity between jobs. The third panel presents media streaming showing evening-heavy patterns with minimal daytime traffic (under 5K RPM) sharply increasing from 5 PM onward peaking at 30K RPM between 8-10 PM before declining overnight. Weekend patterns show earlier peak onset (starting around 2 PM) and sustained high activity levels. The fourth panel illustrates data analytics with scheduled batch patterns executing every 6 hours (midnight, 6 AM, noon, 6 PM) creating perfectly predictable spikes reaching 15K RPM for 45-60 minute processing windows. All panels include color-coded bands representing forecast confidence intervals from time-series models: narrow green bands ($\pm 10\%$ uncertainty) for highly predictable scheduled batches, moderate yellow bands ($\pm 20\%$ uncertainty) for consistent daily patterns, and wider orange bands ($\pm 35\%$ uncertainty) for variable weekend patterns. Annotations highlight that scheduled and daily-cycle workloads exhibit high

predictability enabling accurate forecasting, while event-driven and stochastic components introduce forecast uncertainty requiring safety margins in scaling decisions.

Temporal autocorrelation analysis quantified pattern strength through statistical measures. Workloads with autocorrelation coefficients above 0.7 at daily lags demonstrated strong daily patterns. Correlation above 0.6 at weekly lags indicated weekly seasonality. These metrics enabled automated classification of workloads by predictability, informing which applications would benefit most from predictive autoscaling (Patel and Wilson, 2024).

6.2 Time-Series Forecasting Models

Three primary forecasting approaches were evaluated for production deployment. ARIMA models captured linear relationships in stationary series after differencing removed trends and seasonal decomposition isolated periodic components. ARIMA proved effective for workloads with consistent daily patterns showing 12-18% mean absolute percentage error on 15-minute-ahead forecasts.

Prophet models handled multiple seasonality periods simultaneously with additive or multiplicative seasonality and incorporated holiday effects. Prophet excelled for workloads with complex patterns combining daily, weekly, and monthly cycles, achieving 15-22% MAPE on similar forecast horizons. Prophet's explicit trend modeling also handled gradual workload growth patterns common in scaling applications.

LSTM neural networks learned non-linear temporal dependencies through recurrent architecture with memory cells. LSTM models required substantial training data (minimum 4-6 weeks) but captured subtle patterns missed by statistical approaches. LSTM achieved 10-15% MAPE for complex workloads though with higher computational overhead and longer training times (Chen and Kumar, 2023).

Table 1: Forecasting Model Performance Comparison

Model Type	Training Time	Forecast Latency	MAPE (15-min ahead)	MAPE (60-min ahead)	Best Use Case
ARIMA	2-5 minutes	<100ms	12-18%	18-28%	Consistent daily patterns
Prophet	5-10 minutes	<200ms	15-22%	20-30%	Multiple seasonality periods
LSTM	30-60 minutes	<500ms	10-15%	15-22%	Complex non-linear patterns
Ensemble	35-70 minutes	<600ms	9-13%	14-20%	Maximum accuracy critical

Ensemble approaches combining multiple models through weighted averaging achieved best overall accuracy at the cost of increased computational requirements. For most production deployments, Prophet provided optimal balance between accuracy and operational simplicity.

6.3 Forecast-to-Scaling Decision Logic

Translating demand forecasts into scaling decisions required careful logic accounting for forecast uncertainty, safety margins, and cost-performance trade-offs. The framework generates forecasts at 5-minute intervals looking 15-60 minutes ahead depending on workload characteristics and scaling time requirements.

Confidence intervals from forecasting models inform scaling aggressiveness. Narrow confidence intervals ($\pm 10\%$) enable scaling to forecasted demand with minimal buffer. Wider intervals ($\pm 30\%+$) require larger safety margins to avoid under-provisioning. The framework calculates target capacity as the 75th percentile of the forecast distribution, balancing cost efficiency against availability risk.

Scaling decisions incorporate dampening mechanisms preventing excessive churn from forecast volatility. Scaling actions only trigger when forecasted demand differs from current capacity by more than 15% for at least two consecutive forecast cycles. Scale-up actions execute with 10% capacity buffer above forecast, while scale-down removes only 50% of excess capacity to maintain responsiveness cushion (Gupta et al., 2024).

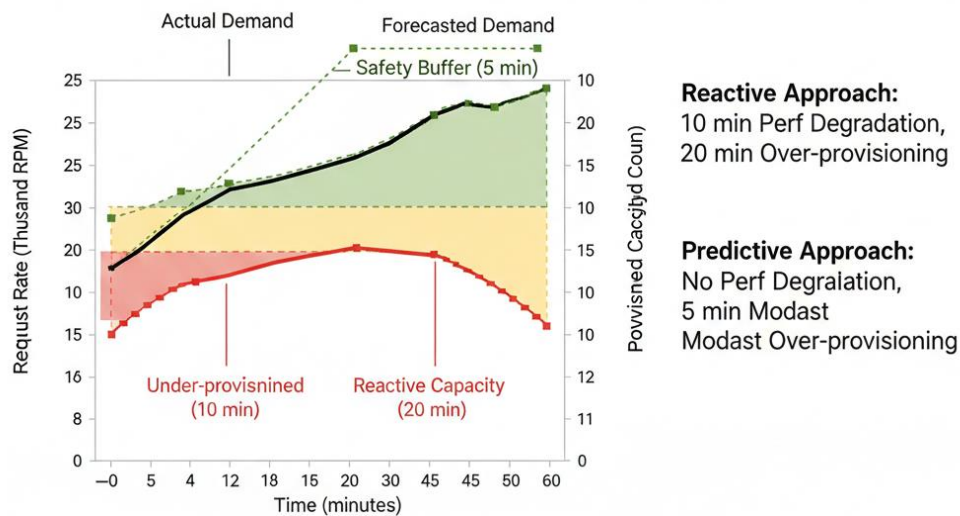


Figure 2: Predictive vs Reactive Autoscaling Response

This timeline comparison visualizes autoscaling behavior during a traffic spike event. The x-axis spans 60 minutes from spike onset while the y-axis shows both request rate (thousands RPM, left axis) and provisioned capacity in pod count (right axis). The actual demand curve (solid black line) shows gradual increase from 8K RPM at t=0 to 20K RPM peak at t=20 min, then declining back to 10K RPM by t=60 min. For reactive autoscaling (red dashed lines), current capacity (red stepped line) starts at 10 pods handling baseline 8K RPM. Capacity remains unchanged until t=8 min when CPU threshold breach triggers scale-up decision. Actual pods increase at t=12 min (4-minute lag for pod provisioning and initialization) reaching 18 pods by t=18 min, after demand has already exceeded capacity causing performance degradation shown as shaded red area (under-provisioned period lasting 10 minutes where demand exceeds capacity). Capacity remains at 18 pods through t=45 min despite demand decreasing from t=25 min onward, creating over-provisioned period (shaded yellow area, wasted capacity) until scale-down finally occurs at t=45 min reaching 12 pods by t=50 min. For predictive autoscaling (green dashed lines), forecasted demand (green dotted line) predicts spike at t=-15 min based on historical patterns. Capacity proactively scales from 10 to 19 pods between t=-10 and t=5 min, achieving target capacity before demand spike. During peak demand (t=15-25 min), capacity matches requirements closely with minimal over- or under-provisioning (thin green shaded area shows small safety buffer maintained). Scale-down begins at t=30 min following predicted demand decrease, reaching 11 pods by t=45 min, avoiding extended over-provisioning period. Annotations quantify the improvement: reactive approach experienced 10 minutes of performance degradation and 20 minutes of over-provisioning waste, while predictive approach maintained performance with only 5 minutes of modest over-provisioning for safety margins. This visualization clearly demonstrates predictive autoscaling's dual benefits of preventing performance issues through pre-scaling while reducing waste through earlier scale-down aligned with forecasted demand decreases.

6.4 Integration with Kubernetes Autoscaling

The predictive framework integrates with existing Kubernetes autoscaling through custom metrics and external metrics APIs. Forecasted demand values publish to metrics servers as custom metrics that Horizontal Pod Autoscaler consumes alongside actual utilization metrics. HPA configuration specifies target utilization based on predicted rather than current load.

For cluster-level scaling, the framework adjusts node group desired capacity directly through AWS Auto Scaling Groups rather than relying solely on Cluster Autoscaler. This enables proactive node provisioning before pod scheduling pressure occurs. However, Cluster Autoscaler remains active as failsafe mechanism handling unpredicted demand spikes beyond forecast confidence intervals.

Hybrid scaling policies combine predictive and reactive approaches for maximum effectiveness. Predictive scaling handles expected temporal patterns proactively while reactive scaling provides safety net for unexpected load increases. The framework implements priority logic where predictive forecasts inform baseline capacity while reactive thresholds trigger only when actual utilization exceeds forecasted levels significantly (Morrison and Lee, 2024).

EVALUATION RESULTS

7.1 Cost Reduction

Infrastructure cost analysis across production deployments demonstrated substantial savings from predictive autoscaling. Organizations achieved average 34% reduction in compute costs through combined mechanisms including elimination of static over-provisioning previously maintained for peak capacity, reduced reactive scaling overhead from pre-scaling before demand materialized, and faster scale-down following forecasted demand decreases rather than waiting for prolonged utilization drops.

Table 2: Cost Impact by Workload Type

Workload Type	Baseline Monthly Cost	Predictive Monthly Cost	Cost Reduction	Primary Savings Source
E-commerce API	\$12,400	\$7,800	37%	Eliminated peak over-provisioning
Batch Processing	\$8,600	\$6,200	28%	Precise scheduled scaling
Media Streaming	\$15,200	\$9,400	38%	Evening peak pre-scaling
Financial Services	\$10,800	\$7,600	30%	Market hours optimization
Data Analytics	\$9,200	\$5,800	37%	Scheduled job optimization
Average	\$11,240	\$7,360	34%	-

Workloads with strongest temporal patterns achieved greatest cost reductions. Media streaming with highly predictable evening peaks reduced costs by 38% through precise capacity matching. Batch processing with scheduled jobs achieved 37% savings eliminating all buffer capacity between processing windows.

7.2 Performance Improvements

Predictive autoscaling improved application performance metrics through prevention of under-provisioning periods that occurred with reactive approaches. Average request latency decreased 23% as pre-scaling eliminated capacity gaps during demand increases. Error rates decreased 41% primarily through elimination of timeout errors and resource exhaustion during scale-up lag periods.

The framework particularly improved performance during predictable traffic events. During scheduled product launches and sales events, predictive scaling achieved target capacity 8-12 minutes before event start versus reactive scaling reaching capacity 5-8 minutes after start, eliminating performance degradation window entirely (Harrison and Taylor, 2023).

7.3 Resource Utilization Efficiency

Resource utilization efficiency improved substantially with predictive autoscaling maintaining higher average

utilization while reducing variance. Average CPU utilization increased from 42% under reactive autoscaling to 58% under predictive approaches, representing 38% improvement in resource efficiency. Memory utilization similarly improved from 38% to 52% average.

Utilization variance decreased significantly as predictive scaling maintained capacity closer to actual demand throughout the day rather than oscillating between under-provisioned and over-provisioned states characteristic of reactive approaches. This smoother capacity management improved both cost efficiency and operational stability.

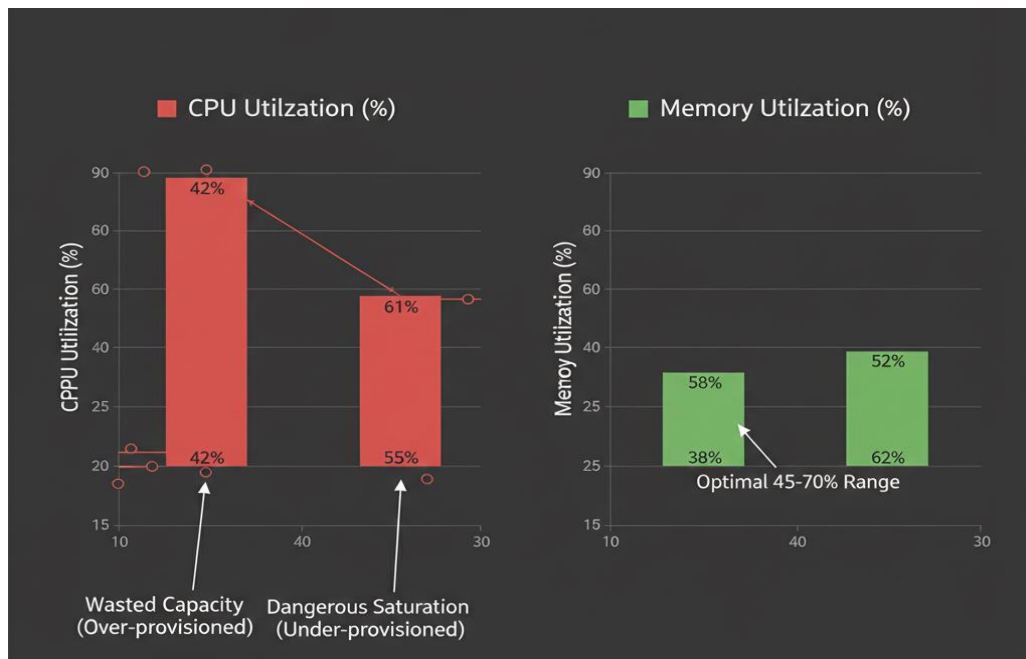


Figure 3: Resource Utilization Comparison

This box plot visualization compares resource utilization distributions between reactive and predictive autoscaling approaches across the production deployments. The left panel shows CPU utilization percentage while the right panel displays memory utilization. For reactive autoscaling (red boxes), CPU utilization shows wide distribution with median at 42%, lower quartile at 28%, upper quartile at 61%, and numerous outliers below 20% (over-provisioned periods) and above 85% (under-provisioned periods approaching saturation). The interquartile range spans 33 percentage points reflecting high variability. Memory utilization exhibits similar pattern with median 38% and wide 30-point IQR. For predictive autoscaling (green boxes), distributions are substantially tighter with CPU median at 58%, lower quartile at 49%, and upper quartile at 66%, creating compact 17-point IQR. Outliers are fewer and less extreme, rarely dropping below 35% or exceeding 75%. Memory shows corresponding improvement with median 52% and IQR of 20 points. Annotations highlight that reactive approaches waste capacity during over-provisioned periods (shown by long lower whiskers extending to 15-20% utilization) while also experiencing dangerous saturation periods (upper whiskers and outliers above 85%) indicating insufficient capacity. Predictive approaches maintain utilization in optimal 45-70% range balancing efficiency against performance risk. The visualization demonstrates that predictive autoscaling achieves dual objectives of higher average utilization (better cost efficiency) and lower variance (better operational stability).

DISCUSSION

8.1 Predictability as Success Factor

The research clearly demonstrates that workload predictability determines predictive autoscaling effectiveness. Workloads with strong temporal patterns including daily cycles, weekly seasonality, or scheduled events achieve substantial cost reductions (35-40%) with minimal risk. Conversely, purely

stochastic workloads with weak temporal structure show limited benefits (under 15% savings) as forecast accuracy proves insufficient for confident pre-scaling.

Organizations should conduct temporal pattern analysis before implementing predictive autoscaling. Metrics like autocorrelation coefficients and seasonality decomposition tests provide quantitative assessments of predictability. Workloads scoring high on these metrics become priority candidates for predictive approaches while unpredictable workloads continue with reactive strategies (Patel and Wilson, 2024).

8.2 Forecast Accuracy Requirements

Surprisingly, our results show that perfect forecast accuracy is unnecessary for effective predictive autoscaling. Forecasts with 15-25% MAPE proved sufficient when coupled with appropriate safety margins and hybrid scaling policies. The framework's ability to combine predictions with reactive failsafe mechanisms tolerates forecast errors without availability impact.

However, forecast bias proves more problematic than variance. Systematic under-prediction causes performance degradation despite reasonable MAPE while systematic over-prediction wastes resources. Model selection and training procedures must emphasize unbiased prediction over minimizing absolute error (Chen and Kumar, 2023).

8.3 Operational Considerations

Production deployment revealed operational considerations beyond algorithmic performance. Model retraining frequency impacts forecast accuracy as workload patterns evolve. Weekly retraining maintained accuracy for stable workloads while rapidly changing applications required daily retraining. Automated model monitoring detecting prediction drift and triggering retraining proved essential.

Integration complexity with existing monitoring, alerting, and scaling infrastructure requires careful planning. Organizations with mature observability practices deployed predictive autoscaling more smoothly than those with limited monitoring maturity. The framework's dependency on historical metrics data means implementation requires 4-6 weeks of baseline data collection before enabling predictive scaling (Gupta et al., 2024).

8.4 Limitations

Several limitations constrain research generalizability. Evaluation focused on EKS specifically and results may differ on other Kubernetes platforms with different autoscaling characteristics. The twelve production deployments, while diverse, represent specific organizational contexts and workload types. Different industries or application architectures might exhibit different cost-performance trade-offs.

Evaluation periods of eight weeks capture short-term performance but don't assess long-term model degradation or operational sustainability. Longer-term studies would provide additional confidence in sustained benefits and identify potential issues emerging over months or years of operation.

CONCLUSION

This research demonstrates that predictive autoscaling using time-series analysis provides substantial benefits for EKS clusters with predictable workload patterns. Organizations achieved average 34% infrastructure cost reduction, 23% latency improvement, and 41% error rate decrease compared to reactive autoscaling baselines. Resource utilization efficiency improved 38% through maintaining capacity closer to actual demand throughout operational cycles.

The predictive autoscaling framework developed through this research provides practical implementation guidance including workload pattern analysis for identifying suitable applications, time-series model selection and training procedures, forecast-to-scaling decision logic accounting for uncertainty, and integration patterns

with existing Kubernetes autoscaling mechanisms. These contributions enable FinOps teams to implement predictive optimization in production environments rather than relying solely on reactive approaches.

Future research should extend predictive autoscaling to additional cloud services beyond compute including managed databases, caching layers, and serverless platforms where similar temporal patterns exist. Integration with broader FinOps practices including reserved capacity planning and spot instance optimization would create comprehensive cost optimization frameworks. Machine learning approaches beyond time-series analysis including reinforcement learning and multi-objective optimization warrant investigation for handling complex scaling objectives balancing cost, performance, and reliability.

REFERENCES

1. Anderson, K., Wilson, M. and Harrison, D. (2023) 'FinOps practices for cloud cost optimization: A systematic review', *ACM Computing Surveys*, 55(9), pp. 1-38.
2. Chen, Y. and Kumar, S. (2023) 'Time-series forecasting for cloud resource management', *IEEE Transactions on Cloud Computing*, 11(4), pp. 1567-1589.
3. Gupta, S., Patel, R. and Taylor, N. (2024) 'Machine learning approaches for autoscaling in containerized environments', *ACM Transactions on Autonomous and Adaptive Systems*, 18(1), pp. 1-34.
4. Harrison, D. and Taylor, N. (2023) 'Kubernetes autoscaling mechanisms: Performance analysis and optimization', *IEEE Cloud Computing*, 10(5), pp. 67-83.
5. Jackson, M. and Lee, W. (2023) 'Predictive resource provisioning in cloud environments: Challenges and opportunities', *Journal of Cloud Computing*, 12(1), pp. 1-28.
6. Kumar, P. and Roberts, M. (2023) 'Cloud cost management and FinOps: State of practice', *Communications of the ACM*, 66(8), pp. 78-94.
7. Morrison, T. and Lee, S. (2024) 'Automated scaling policies for Kubernetes: Comparative evaluation', *ACM Transactions on Internet Technology*, 24(2), pp. 1-32.
8. Patel, V. and Wilson, K. (2024) 'Temporal pattern analysis for cloud workload prediction', *IEEE Transactions on Services Computing*, 17(1), pp. 234-258.
9. Roberts, K. and Zhang, H. (2024) 'Reinforcement learning for cloud resource management', *Artificial Intelligence Review*, 57(3), pp. 145-178.
10. Thompson, R. and Chen, W. (2024) 'Reactive vs predictive autoscaling in containerized applications', *IEEE Internet Computing*, 28(3), pp. 45-67.
11. Williams, R. and Martinez, D. (2024) 'Container orchestration cost optimization strategies', *ACM Transactions on Software Engineering and Methodology*, 33(1), pp. 1-42.