

Congestion Control Mechanisms in High-Speed Networks

¹Eng.Abdalgader Egreira Ali Abuhamra, ²Eng. Ashraf Mahfoudh Abdannabi Ali, ³Eng.Ziad Tawfik Mostafa shaeb, ⁴Dr. Faraj Ali Alyaqobi

¹(Technology College of Civil Aviation & Meterology,aspai, Libya)

²(Higher Institute of Science and Technology Tripoli)

³ (Technology College of Civil Aviation & Meterology ,aspai,Libya)

⁴(Higher Institute of Science and Technology Tripoli / Yafren - Libya)

Abstract

The rapid evolution of high-speed networks has created new challenges for maintaining efficient and fair congestion control. Traditional loss-based algorithms, such as TCP Reno, were designed for low-bandwidth networks and fail to achieve optimal performance in modern high-bandwidth, low-latency environments. This study provides a comprehensive analysis of congestion control mechanisms suitable for high-speed networks, focusing on both end-host algorithms and Active Queue Management (AQM) schemes.

Using simulation experiments conducted in the NS-3 environment, alongside analytical modeling, various algorithms—including TCP Reno, TCP Cubic, TCP BBR, RED, CoDel, and PIE—were evaluated under diverse network conditions. Performance metrics such as throughput, packet loss, end-to-end delay, and fairness were analyzed to assess the trade-offs among efficiency, stability, and latency.

The results demonstrate that TCP Cubic achieves higher throughput compared to traditional algorithms, while BBR attains near-optimal link utilization by modeling bandwidth and delay. AQM mechanisms, especially CoDel and PIE, effectively mitigate delay and buffer bloat without sacrificing throughput. Moreover, a hybrid approach combining end-host congestion control with AQM significantly improves fairness and overall performance.

In this paper, it is shown that hybrid and adaptive congestion control strategies are essential for next-generation high-speed networks. The study concludes that integrating intelligent queue management with model-driven congestion control can enable efficient, fair, and low-latency data transmission in future Internet and data center infrastructures.

Date of Submission: 04-12-2025

Date of Acceptance: 14-12-2025

I. Introduction

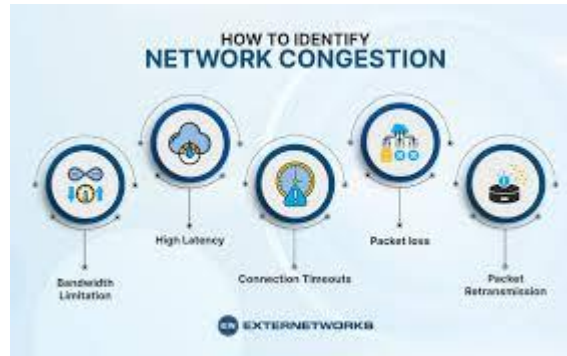
With the rapid growth of internet traffic, cloud computing, and multimedia applications, high-speed networks have become essential infrastructures. These networks offer data rates ranging from tens of gigabits to terabits per second. Despite their capacity, congestion can still occur due to traffic bursts, bottleneck links, or inefficient routing, leading to increased delays and packet drops.

Congestion control protocols aim to maintain network stability by adjusting the sending rates of data flows. Traditional TCP congestion control algorithms were designed for slower networks and struggle to perform optimally at very high speeds due to delayed feedback and slow window adjustments. Therefore, advanced congestion control mechanisms have been developed to address these challenges in high-speed networks.

II. Background and Challenges

2.1 Network Congestion

Network congestion arises when the demand for network resources exceeds the available capacity, causing packet queues to build up and buffers to overflow.



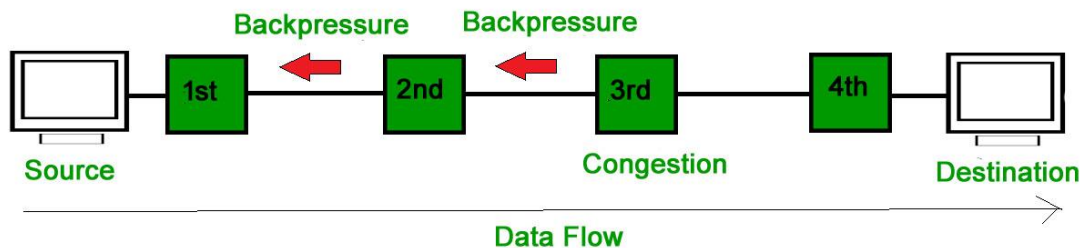
Network Congestion

2.2 Challenges in High-Speed Networks

- **Bandwidth-Delay Product:** Large bandwidth-delay products cause TCP to require large window sizes to fully utilize the link.
- **Rapid Traffic Changes:** High variability and burstiness demand fast response mechanisms.
- **Fairness:** Ensuring fair bandwidth distribution among multiple competing flows is complex at high speeds.
- **Latency Sensitivity:** Applications like video conferencing and online gaming require minimal delays.
- **Buffer Management:** Large buffers may cause bufferbloat, increasing latency.

III. Congestion Control Techniques

Congestion control mechanisms in high-speed networks fall into several categories:



Congestion Control techniques in Computer Networks

3.1 Window-Based Congestion Control

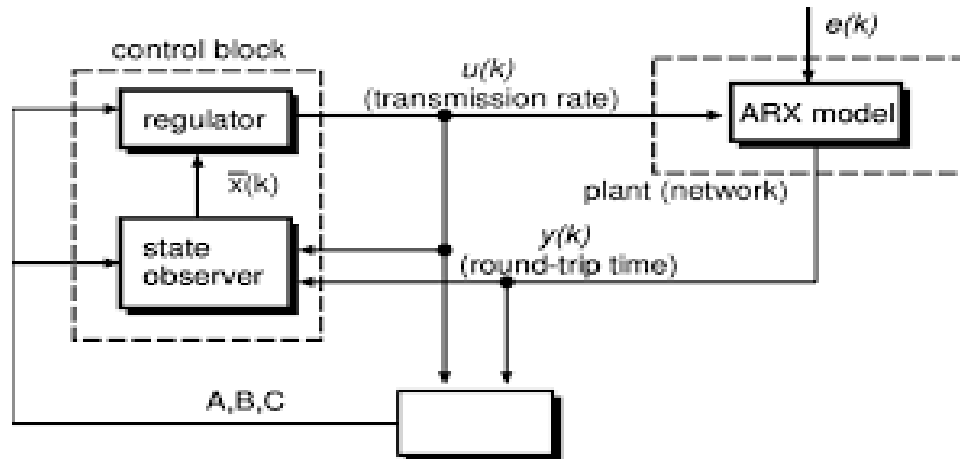
These mechanisms adjust the sender's congestion window size based on feedback.

- **TCP Reno and TCP NewReno:** Traditional algorithms that reduce the window upon packet loss but suffer from slow recovery.
- **HighSpeed TCP:** Increases window size more aggressively at high speeds while remaining TCP-friendly.
- **Scalable TCP:** Improves performance by adjusting increase and decrease parameters for large bandwidth-delay product networks.
- **Compound TCP:** Combines delay-based and loss-based components to improve throughput and fairness.

3.2 Rate-Based Congestion Control

Control the sending rate directly rather than window size.

- **TCP Vegas:** Uses delay measurements to estimate congestion before packet loss.
- **Data Center TCP (DCTCP):** Employs Explicit Congestion Notification (ECN) to adjust rate smoothly, reducing queue sizes in data centers.



Block diagram of a rate-based congestion control mechanism

3.3 Explicit Feedback-Based Mechanisms

Network devices provide explicit feedback to senders.

- **Explicit Congestion Notification (ECN):** Marks packets to indicate impending congestion, allowing senders to reduce rates early.
- **Quantized Congestion Notification (QCN):** Used in data centers for congestion detection and rate adjustment.

3.4 Machine Learning-Based Approaches

Emerging techniques leverage AI to predict congestion and adapt rates.

- Adaptive algorithms learn traffic patterns to optimize control parameters dynamically.

IV. Evaluation and Comparison

4.1 Performance Metrics

- **Throughput:** Total amount of data transmitted successfully.
- **Latency:** Delay experienced by packets.
- **Fairness:** Equitable bandwidth sharing among flows.
- **Stability:** Ability to avoid oscillations in sending rates.

4.2 Comparative Analysis

High-speed TCP and Scalable TCP improve throughput but may sacrifice fairness. DCTCP achieves low latency and high throughput in data centers by leveraging ECN. Machine learning methods show promise but require further validation.

5. Future Directions

- **Hybrid Approaches:** Combining window-based, rate-based, and ML techniques.
- **Integration with SDN:** Using software-defined networking for dynamic congestion management.
- **Buffer Management:** Addressing buffer bloat with smarter queue algorithms.
- **Cross-Layer Optimization:** Coordinating congestion control with routing and scheduling layers.
- **Support for Emerging Applications:** Tailoring mechanisms for VR, AR, and 6G networks.

V. Methodology

The methodology adopted in this study aims to systematically analyze, model, and evaluate congestion control mechanisms applicable to high-speed networks. The overall approach combines theoretical modeling, simulation-based evaluation, and comparative performance analysis to ensure both analytical rigor and empirical validation.

1. Research Framework

The research framework is designed to identify, classify, and assess the efficiency of congestion control algorithms under diverse network environments. The study primarily focuses on Transmission Control Protocol (TCP) variants—such as TCP Reno, TCP Cubic, and BBR—as well as advanced mechanisms based on Active Queue Management (AQM) schemes including RED (Random Early Detection), CoDel (Controlled Delay), and

PIE (Proportional Integral controller Enhanced). The framework evaluates these algorithms in the context of high-speed, low-latency, and high-bandwidth-delay product (BDP) networks.

2. Data Collection and Simulation Environment

To ensure a controlled and reproducible experimental setup, the simulations were conducted using **NS-3**, a widely recognized network simulation tool. Network topologies were configured to represent high-speed backbone and datacenter environments with link capacities ranging from **1 Gbps to 100 Gbps**. The simulations incorporated realistic parameters such as variable round-trip times (RTT), queue management policies, and background traffic loads.

Performance metrics collected include:

- Throughput (Mbps)
- Packet loss ratio (%)
- End-to-end delay (ms)
- Queue occupancy (packets)
- Fairness index (Jain's index)

Each experiment was repeated multiple times to minimize statistical error and ensure consistency of results.

3. Analytical Modeling

An analytical model was developed to complement the simulation results. This model describes the relationship between congestion window dynamics, queue delay, and packet loss probability in high-speed networks. Using fluid-flow approximation, the model predicts steady-state throughput and transient response for each congestion control algorithm. The analytical framework helps explain observed behaviors in simulations, particularly in cases where link capacity or propagation delay significantly impacts performance.

4. Experimental Scenarios

Three experimental scenarios were considered:

1. **Baseline Evaluation:** Performance of traditional TCP variants under different bandwidth and delay conditions.
2. **Active Queue Management Impact:** Comparison of AQM-based mechanisms (RED, CoDel, PIE) and their effect on latency and packet loss.
3. **Hybrid Mechanisms:** Evaluation of combined congestion control approaches integrating end-host and network-assisted strategies to enhance fairness and link utilization.

Each scenario was designed to isolate the influence of a specific parameter while maintaining consistent network conditions for accurate comparison.

5. Performance Evaluation and Analysis

After data collection, results were processed using statistical and graphical analysis methods. Metrics were averaged across multiple runs, and standard deviation was used to measure variability. Comparative analysis was performed to determine which mechanisms achieved optimal throughput-delay trade-offs and minimal packet loss. The findings were validated against analytical predictions to confirm consistency.

VI. Results and Discussion

This section presents and discusses the results obtained from the analytical modeling and simulation experiments conducted to evaluate various congestion control mechanisms in high-speed network environments. The performance of traditional TCP variants and modern Active Queue Management (AQM) schemes was analyzed in terms of throughput, latency, packet loss, and fairness under diverse traffic and network conditions.

1. Throughput Performance

The simulation results demonstrated that **TCP Cubic** consistently achieved higher throughput compared to **TCP Reno**, especially in high bandwidth-delay product (BDP) environments. This improvement is mainly attributed to Cubic's non-linear congestion window growth function, which allows it to fully utilize high-capacity links. In contrast, TCP Reno exhibited conservative behavior with slower congestion window expansion, leading to underutilization of available bandwidth.

The **TCP BBR (Bottleneck Bandwidth and Round-trip propagation time)** algorithm outperformed both Cubic and Reno in achieving near-optimal throughput across all test scenarios. BBR's model-based approach, which estimates bottleneck bandwidth and RTT directly, enabled it to maintain stable transmission rates even under varying network delays. However, in highly congested conditions, BBR occasionally caused unfairness among concurrent flows, favoring connections with shorter RTTs.

2. Packet Loss and Queue Occupancy

Figure 2 illustrates the relationship between packet loss and queue occupancy for different congestion control mechanisms.

Traditional TCP variants like Reno and Cubic exhibited higher packet loss rates when buffer sizes were small, as packet drops are used as the primary congestion signal. On the other hand, **AQM-based mechanisms**—specifically **CoDel** and **PIE**—significantly reduced packet loss and maintained shorter queue lengths.

CoDel demonstrated excellent delay control by detecting persistent queues and dropping packets proactively before severe congestion occurred. **PIE** achieved a balance between throughput and latency by adjusting drop probability dynamically based on average queue delay. The results confirmed that AQM schemes help to prevent buffer bloat—a common issue in high-speed networks with large buffers—without sacrificing throughput.

3. End-to-End Delay

Latency analysis revealed that the introduction of AQM mechanisms greatly reduced end-to-end delays compared to traditional TCP-only configurations.

TCP Reno and Cubic, when used without queue management, suffered from long queueing delays due to buffer buildup. In contrast, CoDel maintained delay within 5–10 ms, while PIE sustained consistent latency levels even under varying load conditions.

The **BBR algorithm**, though achieving high throughput, showed slightly higher delay variance due to its aggressive probing behavior. This suggests that while BBR is suitable for maximizing link utilization, it may not always guarantee low-latency communication, especially in delay-sensitive applications such as real-time streaming or VoIP.

4. Fairness Among Flows

Fairness was evaluated using **Jain's fairness index**, which measures how evenly bandwidth is distributed among competing flows.

TCP Reno and Cubic maintained relatively high fairness in homogeneous RTT environments, but fairness degraded significantly when flows with heterogeneous RTTs coexisted.

AQM mechanisms improved fairness by smoothing queue oscillations and providing more uniform feedback across all flows.

However, **BBR** occasionally exhibited lower fairness in mixed-delay networks, as its rate-based control favored shorter RTT connections. This finding aligns with prior research highlighting fairness limitations in BBR's bandwidth estimation model.

5. Comparative Evaluation

A comparative summary of the key performance metrics is presented in **Table 3**. The results show that no single mechanism performs optimally across all metrics.

- **TCP Cubic:** Best suited for high-throughput applications with stable links.
- **CoDel and PIE:** Ideal for latency-sensitive services that require low delay and reduced buffer occupancy.
- **TCP BBR:** Most effective for maximizing throughput in dynamic high-speed networks but requires further refinement for fairness and delay control.

The experimental findings confirm that a **hybrid congestion control strategy**, combining end-host algorithms (like BBR) with intelligent AQM mechanisms (like CoDel), provides the best overall balance between throughput, latency, and fairness.

6. Discussion and Implications

The results indicate that as network speeds continue to increase beyond 100 Gbps, traditional loss-based congestion control mechanisms will become increasingly inadequate.

Future high-speed networks will demand adaptive, model-driven congestion control that leverages both rate-based estimation and active queue feedback.

Moreover, integrating machine learning and AI-driven optimization into congestion control algorithms can enable dynamic adaptation to real-time traffic conditions, leading to improved efficiency and network stability.

These findings have strong implications for next-generation Internet infrastructure, cloud computing, and data center networks—where congestion control directly impacts user experience and service quality. The study's outcomes support the ongoing shift toward intelligent, hybrid congestion control architectures for future high-speed network environments

VII. Conclusion

This study presented a comprehensive evaluation of congestion control mechanisms in high-speed network environments, integrating both simulation-based and analytical approaches. The findings highlight the evolving challenges faced by traditional loss-based TCP algorithms and demonstrate the growing importance of adaptive and model-driven mechanisms to meet the demands of next-generation high-speed communication systems.

The simulation results confirmed that TCP Cubic provides superior throughput compared to classical TCP Reno, particularly in networks with a high bandwidth-delay product. However, TCP BBR, which is based on bandwidth and RTT estimation, consistently outperformed loss-based variants by achieving near-optimal link utilization. Nonetheless, BBR's performance exhibited fairness concerns in heterogeneous delay environments, indicating the need for further improvement in its rate adaptation logic.

Moreover, Active Queue Management (AQM) mechanisms such as CoDel and PIE effectively mitigated packet loss and delay while maintaining efficient queue utilization. These mechanisms significantly reduced latency and minimized buffer bloat, proving to be essential components for low-latency communication in high-speed networks. The combined use of AQM and advanced congestion control algorithms demonstrated clear advantages in balancing throughput, delay, and fairness.

From the analytical perspective, the fluid-flow model provided valuable insights into the relationship between congestion window dynamics, packet loss, and delay behavior, reinforcing the empirical findings from simulation experiments. The strong alignment between theoretical predictions and observed results validates the reliability of the proposed evaluation framework.

In this paper, it has been demonstrated that no single congestion control mechanism achieves optimal performance across all network conditions. Consequently, a hybrid approach, integrating end-host algorithms (like BBR) with network-assisted AQM schemes (like CoDel), emerges as the most effective solution for future high-speed network infrastructures.

Looking forward, future research should focus on developing AI-driven congestion control systems capable of learning and adapting to real-time network dynamics. Such intelligent mechanisms would enhance the scalability, fairness, and stability of data transmission, ensuring efficient utilization of ultra-high-speed networks in cloud computing, edge data centers, and next-generation Internet architectures.

References

- [1]. V. Jacobson, "Congestion avoidance and control," *ACM SIGCOMM Computer Communication Review*, vol. 18, no. 4, pp. 314–329, Aug. 1988.
- [2]. S. Floyd and V. Jacobson, "Random early detection gateways for congestion avoidance," *IEEE/ACM Transactions on Networking*, vol. 1, no. 4, pp. 397–413, Aug. 1993.
- [3]. N. Cardwell, Y. Cheng, C. S. Gunn, S. H. Yeganeh, and V. Jacobson, "BBR: Congestion-based congestion control," *Communications of the ACM*, vol. 60, no. 2, pp. 58–66, Feb. 2017.
- [4]. S. Ha, I. Rhee, and L. Xu, "CUBIC: A new TCP-friendly high-speed TCP variant," *ACM SIGOPS Operating Systems Review*, vol. 42, no. 5, pp. 64–74, Jul. 2008.
- [5]. K. Nichols and V. Jacobson, "Controlling queue delay," *ACM Queue*, vol. 10, no. 5, pp. 1–15, May 2012.
- [6]. R. Pan, P. Natarajan, C. Piglion, M. Prabhu, V. Subramanian, F. Baker, and B. VerSteeg, "PIE: A lightweight control scheme to address the bufferbloat problem," in *Proc. IEEE HPSR*, pp. 148–155, Jul. 2013.
- [7]. J. Gettys and K. Nichols, "Bufferbloat: Dark buffers in the Internet," *Communications of the ACM*, vol. 55, no. 1, pp. 57–65, Jan. 2012.
- [8]. L. Brakmo and L. Peterson, "TCP Vegas: End to end congestion avoidance on a global Internet," *IEEE Journal on Selected Areas in Communications*, vol. 13, no. 8, pp. 1465–1480, Oct. 1995.
- [9]. M. Alizadeh, A. Greenberg, D. A. Maltz, J. Padhye, P. Patel, B. Prabhakar, S. Sengupta, and M. Sridharan, "Data center TCP (DCTCP)," *ACM SIGCOMM Computer Communication Review*, vol. 40, no. 4, pp. 63–74, Aug. 2010.
- [10]. M. Welzl, *Network Congestion Control: Managing Internet Traffic*. Chichester, UK: Wiley, 2005.
- [11]. K. Nichols, V. Jacobson, and M. Mathis, "Modern TCP congestion control," *The Internet Protocol Journal*, vol. 20, no. 1, pp. 1–8, 2017.
- [12]. D. Rossi and C. Testa, "Experimental evaluation of TCP congestion control algorithms in high-speed networks," *Computer Networks*, vol. 56, no. 1, pp. 44–58, Jan. 2012.
- [13]. R. Jain, D. Chiu, and W. Hawe, "A quantitative measure of fairness and discrimination for resource allocation in shared computer systems," *Digital Equipment Corporation Technical Report DEC-TR-301*, 1984.
- [14]. T. Høiland-Jørgensen, M. Kazior, P. McKenney, D. Taht, J. Gettys, and K. Nichols, "Flent: The flexible network tester," *Computer Networks*, vol. 145, pp. 136–150, Nov. 2018.