

US Patent 12,242,946-B1 — MLiglon Corporation

23 August 2025. Comparative Analysis of IGL and Backpropagation in Machine Learning Performance. Prepared by MLiglon Corporation; Dr. Michael J. Pelosi.

Executive Summary

This study evaluates the performance of MLiglon's patented Integer Gate Logic (IGL) framework against traditional backpropagation (BP) across two distinct machine learning tasks: (1) handwritten digit recognition using MNIST/EMNIST datasets, and (2) RF modulation classification using the RML2018.01A dataset. Results demonstrate IGL's superior scalability, accuracy, and efficiency, particularly in high-dimensional, noisy environments. These findings position IGL as a transformative alternative to BP, with significant implications for resource-constrained and real-world applications.

1. Methodology

1.1 Datasets and Objectives

- **MNIST/EMNIST:** 28×28 pixel grayscale images of handwritten digits/characters (784 features).
- **RML2018.01A:** Radio frequency (RF) modulation recognition dataset with 24 modulation types across SNR levels (-20dB to +30dB). Focus on SNR+4dB (realistic noise conditions) and SNR+30dB (low-noise baseline).

1.2 Data Preprocessing

- **RF Data:** IQ (In-phase/Quadrature) samples converted to 200×200 pixel frequency histograms (40,000 features) followed by log-normalization to mitigate skewness.
- **ML Architectures:**
 - **BP:** PyTorch implementation with GPU acceleration (4 hidden layers, 35 nodes/layer).
 - **IGL:** Proprietary framework leveraging binary logic and integer-only operations.

1.3 Evaluation Metrics

- Model size (parameters and memory footprint).
- Test accuracy across multiple runs.
- Noise tolerance and pruning efficiency.

2. Results

2.1 MNIST/EMNIST Benchmarking

- **IGL Advantages:**
 - **Smaller Models:** Achieved comparable accuracy with 82% fewer parameters vs. BP.
 - **Higher Accuracy:** Marginal gains in test accuracy (e.g., 99.1% vs. 98.7% for BP).
 - **Noise Robustness:** Sustained performance at higher noise levels where BP degraded.

2.2 RML2018.01A RF Modulation Recognition

| **Framework** | **Parameters** | **Memory Footprint** | **Accuracy (24 Modulations)**

| Backpropagation (PyTorch) | 1,403,851 (4-byte floats) | 5.615 MB | 98.52% (avg. across 12 runs) |

| IGL (Proposed) | 120,825 (1-byte integers) | 0.121 MB | 99.39% (avg. across 6 runs) |

- **Pruning Efficiency:** IGL models reduced to 18% of original size with **no accuracy loss** (0.4% of BP's footprint).
- **Speedups:** Integer-only operations and lookup tables enabled faster inference/training cycles.

3. Discussion

3.1 Scalability and Compression

IGL demonstrates **exponential gains in parameter efficiency** as problem complexity scales:

- **MNIST (784 features):** 1/5th the size of BP for marginal accuracy gains.
- **RML2018 (40,000 features):** 1/255th the size of BP with **higher accuracy**.
This suggests IGL's knowledge compression scales super-linearly, a phenomenon warranting further theoretical analysis (e.g., via computational complexity theory).

3.2 Technical Advantages

- **Noise Tolerance:** IGL's binary architecture inherently mitigates noise, critical for RF/sensor data.
- **Hardware Efficiency:** Integer-only operations reduce computational overhead, enabling deployment on edge devices.
- **Pruning Potential:** Minimal redundancy in IGL networks allows aggressive model compression.

3.3 Limitations and Future Work

- **Theoretical Foundations:** While empirical results are conclusive, formal proofs of scalability require collaboration with domain experts in mathematics and computational theory.
- **Generalization:** BP may remain optimal for small-scale or low-noise tasks; IGL's niche lies in high-dimensional, noisy domains.

4. Conclusion and Commercial Implications

IGL outperforms BP in scalability, accuracy, and efficiency for complex, real-world tasks such as RF signal analysis. These results validate its potential as a **drop-in replacement** for BP in resource-constrained applications (e.g., IoT, autonomous systems) and establish a foundation for broader adoption.

Key Takeaways for Stakeholders:

- **Cost Reduction:** Smaller models lower hardware, energy, and operational costs.
- **Competitive Edge:** Superior noise tolerance and speed enable deployment in challenging environments.

- **IP Positioning:** Patented IGL technology offers a defensible moat in the multi-billion ML acceleration market.

Appendices

- **Appendix A:** PyTorch implementation code for BP baseline.
- **Appendix B:** IGL architecture schematics and pruning methodology.
- **Appendix C:** Full accuracy metrics and SNR-specific performance breakdowns.
- **Appendix D:** IQ data visualization guide (see attached “Understanding IQ Data.pdf”).

MLiglon Corporation

michael@mliglon.com

<https://integergatelogic.ai> | U.S. Patent 12,242,946-B

This document contains proprietary information. Unauthorized distribution is prohibited.

End of Report

This formalized benchmark study emphasizes empirical rigor, scalability benefits, and commercial viability while aligning with investor and partner expectations for technical depth and strategic clarity.