

The Moganti-Nyapathi Limit™: A Framework for Quantifying Reliable Intelligence Efficiency in Artificial Intelligence Systems

Madhav Moganti, Moganti Labs LLC, mmoganti@mogantilabs.com

and

Ramana Nyapathi, Aganitha Cognitive Services, ramana@aganitha.ai

July 2026

Abstract

Artificial intelligence has advanced rapidly through scaling of model parameters, datasets, "and computational resources. Scaling laws have demonstrated predictable improvements in model "capability; however, capability alone does not fully describe the value of an AI system deployed "in real-world environments.

This paper introduces the *Moganti-Nyapathi Limit™*, a proposed framework for measuring Reliable "Intelligence Efficiency (RIE): the amount of trustworthy intelligence generated per unit of "energy, information quality, and verification effort.

The framework introduces the *Moganti-Nyapathi Intelligence Efficiency Index (MNIET)*, which "combines capability, accuracy, robustness, and trustworthiness with resource constraints. "The framework integrates concepts from thermodynamics of computation, information theory, AI "scaling laws, synthetic-data stability, and verification science.

The *Moganti-Nyapathi Limit™* is not proposed as an established physical law. Instead, it is proposed "as a falsifiable research hypothesis: future AI progress should increasingly be measured by the "efficiency with which reliable intelligence is produced rather than by model scale alone.

Keywords

Moganti-Nyapathi Limit™; AI scaling; reliable intelligence; intelligence efficiency; AI economics; thermodynamics of computation; information theory; trustworthy AI.

1. Introduction

The dominant paradigm of modern artificial intelligence has been scaling. Larger neural "networks, increased datasets, and greater computational resources have produced substantial "advances across language, vision, and multimodal intelligence.

However, deployment introduces constraints that are not captured by benchmark accuracy alone. "AI systems must be energy efficient, robust under changing conditions, trained on reliable "information, and sufficiently verifiable for human use.

This motivates a change in the central question of AI progress:

Traditional question:

How large is the AI model?

Proposed question:

How much reliable intelligence do we obtain per unit of energy, data quality, and verification effort?

The Moganti-Nyapathi Limit™ provides a framework for studying this transition.

2. Theoretical Foundation

The framework builds upon four established areas.

Thermodynamics of computation establishes that computation is a physical process with energy "constraints. Information theory shows that intelligence depends on extracting task-relevant "representations rather than simply accumulating information. Scaling-law research describes "capability improvements from increased compute and data. Verification research examines the "difficulty of ensuring complex AI systems behave reliably.

The contribution of this work is integrating these perspectives into a unified efficiency metric.

3. The Moganti-Nyapathi Limit™

The Moganti-Nyapathi Limit™ describes an efficiency frontier: the maximum reliable intelligence "obtainable from available physical, informational, and verification resources.

The central hypothesis is:

For a specified AI domain and reliability requirement, increasing model scale alone does not "guarantee increasing intelligence efficiency.

Beyond certain regimes, additional scale may require disproportionate increases in energy, data "requirements, and verification complexity.

4. Mathematical Formulation

The Moganti-Nyapathi Intelligence Efficiency Index (*MNIEI*) is defined as:

$$MNIEI = I_R / (E^\alpha \times D^\beta \times V^\gamma)$$

where:

- I_R = Reliable Intelligence Output
- E = Energy consumption
- D = Data resource complexity
- V = Verification effort

The reliable intelligence term is:

$$I_R = \text{Capability} \times \text{Accuracy} \times \text{Robustness} \times \text{Trustworthiness}$$

The parameters α , β , and γ represent application-dependent weighting factors.

Table 1. *MNIEI* Evaluation Framework

Dimension	Measurement	Example Metrics	Purpose
Capability	Task performance	Accuracy, completion rate	Useful intelligence
Energy	Physical resources	Training and inference cost	Efficiency
Data	Information quality	Lineage, diversity	Learning efficiency
Verification	Reliability controls	Testing, monitoring	Trustworthiness

5. Benchmark Methodology

MNIEI evaluation requires normalization of multiple dimensions. A practical benchmark should "measure capability, resource efficiency, reliability, and verification burden.

The objective is not to replace traditional AI benchmarks but to augment them with efficiency "and trust measurements.

Table 2. Evolution of AI Optimization Objectives

Era	Optimization Goal	Primary Limitation
Early AI	Accuracy	Limited data and compute
Deep Learning	Scale	Resource consumption
Foundation Models	Capability	Verification complexity
Future AI	Reliable intelligence efficiency	Need for new metrics

6. Experimental Validation Strategy

A complete MNIEI benchmark requires publicly reproducible measurements.

Candidate measurements include:

- benchmark capability;
- training and inference energy;
- data quality indicators;
- robustness metrics;
- verification requirements.

Because commercial frontier models do not disclose all internal resource information, early "validation should focus on systems with sufficient transparency.

7. Industry and Economic Implications

The framework suggests that future AI competition may shift from model size toward intelligence "efficiency.

Potential industry metrics include:

- intelligence per watt;
- intelligence per dollar of compute;
- intelligence per training example;
- verification efficiency;
- deployment reliability.

For investors, the central question becomes:

Who converts computational investment into reliable economic intelligence most efficiently?

8. Research Agenda

Future research should establish:

1. Standardized *MNIEI* benchmarks.
2. Public AI efficiency databases.
3. Statistical uncertainty methods.
4. Domain-specific weighting functions.
5. Longitudinal studies of AI efficiency scaling.

The goal is to determine whether reliable intelligence efficiency exhibits predictable scientific behavior.

9. Limitations

The Moganti-Nyapathi Limit™ is a proposed framework rather than an established physical law.

Important open questions include:

- universal intelligence measurement;
- normalization across domains;
- weighting parameter selection;
- empirical validation of efficiency frontiers.

These questions define the future research program.

10. Conclusion

The first era of AI was defined by scaling capability.

The next era may be defined by scaling reliable intelligence.

The Moganti-Nyapathi Limit proposes that the ultimate measure of AI progress is not merely the "size of artificial minds, but the efficiency with which those systems transform energy, "information, and computation into trustworthy intelligence.

References

- [1] Kaplan et al., Scaling Laws for Neural Language Models, OpenAI, 2020.
- [2] Hoffmann et al., Training Compute-Optimal Large Language Models, DeepMind, 2022.
- [3] Landauer, R., Irreversibility and Heat Generation in the Computing Process, IBM Journal of Research and Development, 1961.
- [4] Tishby, N., Pereira, F., Bialek, W., The Information Bottleneck Method, 1999.
- [5] Sutton, R., The Bitter Lesson, 2019.
- [6] Amodei et al., Concrete Problems in AI Safety, 2016.
- [7] Russell, S., Human Compatible: Artificial Intelligence and the Problem of Control, 2019.
- [8] Bender et al., On the Dangers of Stochastic Parrots, FAccT, 2021.