

When Empirical Scaling Meets Deterministic Theory

How Meta's ScaleRL Validates the Deterministic Envelope of AI Governance

October 2025 | FERZ LLC

Keywords: Deterministic AI Governance, ScaleRL, LASO(f), Reinforcement Learning Scaling Laws, FERZ Deterministic Envelope

Introduction

Artificial intelligence is entering a new era of predictability. With the publication of *The Art of Scaling Reinforcement Learning Compute for LLMs* (Meta AI & University of Texas, October 2025), reinforcement learning has gained its first reproducible scaling law—validated across 400,000 GPU-hours of systematic experiments and successfully extrapolated to 100,000 GPU-hour training runs with predictive accuracy. The paper demonstrates that RL performance follows a sigmoidal relationship with compute—a law of saturation, not endless growth.

At FERZ, this result resonates deeply. It provides the first large-scale empirical confirmation of a principle that our architecture has embodied from the start: that intelligence, however complex, evolves within deterministic boundaries.

The Deterministic Envelope

The **Deterministic Envelope** is FERZ's internal term for the mathematically bounded region within which an intelligent system behaves predictably and auditably. It unifies the core invariants first formalized in LASO(f):

- **Ceiling (A)**: the maximum admissible fidelity or semantic accuracy
- **Efficiency (B)**: the rate of convergence under finite resources
- **Equilibrium (C_mid)**: the stability point where deterministic behavior stabilizes

While the phrase "deterministic envelope" originated in seismology, FERZ repurposes it here in a completely non-seismic sense: no tremors, just steady governance curves.

Within this envelope, all learning and reasoning operations remain reproducible, bounded by formal constraints that can be mathematically proven and externally audited.

Empirical Parallel: The ScaleRL Curve

Meta's ScaleRL framework expresses reinforcement learning reward growth as:

Empirical Parallel: The ScaleRL Curve

ScaleRL models reward growth as:

$$R_C - R_0 = (A - R_0) \frac{1}{1 + (C_{\text{mid}}/C)^B}$$

where:

- R_C = expected reward at compute C
- R_0 = baseline reward
- A = asymptotic performance ceiling
- B = compute-efficiency exponent (rate of convergence)
- C_{mid} = equilibrium midpoint (inflection point)

These parameters correspond exactly to FERZ's invariants—**Ceiling (A), Efficiency (B), Equilibrium (C_{mid})**.

Meta's experiments demonstrate that this curve remains predictive across:

- Model scales (8B to 17B×16 MoE)
- Context lengths (14K to 32K tokens)
- Batch sizes (768 to 2048)
- Task mixtures (math, code)

This cross-dimensional stability mirrors FERZ's principle that deterministic bounds are **invariant properties of the governance architecture**, not artifacts of specific configurations.

ScaleRL's data confirm that as compute increases, learning follows a deterministic curvature—not an open-ended exponential. Each experiment converges toward a measurable asymptote, validating the bounded-optimization view that FERZ has treated as foundational.

From Architecture to Observation

In LASO(f), DELIA, and STRATA-G, these invariants are operationalized through:

- *Bounded optimization manifolds* that confine semantic transformation

- *Microkernel invariants* enforcing resource predictability
- *Recursive governance layers* ensuring stability before self-modification

ScaleRL now provides empirical evidence that these architectural constraints reflect an underlying law of learning itself. The science has arrived at the same conclusion—not to rival the architecture, but to corroborate it from an independent empirical foundation.

The Convergence of Theory and Practice

What makes this convergence significant is not merely that both frameworks use similar parameters—many systems do. Rather, it's that:

1. Independent derivation: Meta derived these bounds empirically through massive-scale experimentation; FERZ derived them axiomatically from governance requirements.

2. Functional equivalence: Both frameworks treat A, B, and C_{mid} not as fitting parameters but as fundamental properties of the learning process.

3. Predictive power: Both enable extrapolation from smaller-scale experiments to larger deployments—ScaleRL through curve fitting, FERZ through formal constraint propagation.

This is the hallmark of scientific maturity: when independently developed approaches converge on the same mathematical structure.

Implications

Scientific: Deterministic bounds are no longer theoretical; they are measurable and reproducible across diverse experimental conditions.

Technical: Governance frameworks can use scaling parameters to predict safe operating limits—including asymptotic performance ceilings and compute-to-capability ratios—enabling resource-aware alignment strategies.

Regulatory: Auditable invariants provide a path toward enforceable AI compliance, with mathematical proofs that can be verified by third parties.

Philosophical: Intelligence is not unbounded emergence but bounded convergence—a principle that reconciles capability advancement with safety guarantees.

Conclusion

ScaleRL and FERZ converge on the same insight from opposite directions: one from data, the other from design. The result is a unifying law of predictable intelligence—the **Deterministic Envelope**—where learning remains measurable, governable, and aligned by construction.

As the AI field celebrates new scaling laws, FERZ continues to focus on the next frontier: **provable governance**—knowing the limits, governing the curvature, and ensuring that intelligence remains deterministically aligned with human intent.

The question is no longer *whether* AI systems have deterministic bounds, but **how precisely we can characterize and enforce them**.

Reference

Khatri, D., Madaan, L., Tiwari, R., Bansal, R., Duvvuri, S.S., Zaheer, M., Dhillon, I.S., Brandfonbrener, D., & Agarwal, R. (2025). *The Art of Scaling Reinforcement Learning Compute for LLMs*. arXiv:2510.13786v1 [cs.LG]

About FERZ LLC

FERZ LLC develops deterministic governance frameworks for artificial intelligence systems, including LASO(f), DELIA, and STRATA-G. Our architectures ensure that AI capabilities remain mathematically bounded, formally verifiable, and aligned with institutional requirements.

Contact: contact@ferzconsulting.com | www.ferzconsulting.com

© 2025 FERZ LLC. All rights reserved.