

# Unified Adaptation Theorem: Convergence of Composed Adaptive Systems via Interaction Matrices and Higher-Order Convergence Diagnostics

Rod Higgins  
Senuamedia

March 2026

## Abstract

We present a mathematical framework for proving convergence of systems composed of multiple adaptive subsystems operating on shared parameters. The central result, the Unified Adaptation Theorem, provides sufficient conditions under which such composed systems converge to a well-defined invariant set. The framework introduces cosine-scaled gradient projection (100% conflict resolution), normalised Lyapunov functions, learnable interaction matrices, higher-order convergence scores, the I-ratio equilibrium criterion ( $I = -1/2$ ), B-flow precision refinement, and desire as Bayesian regularisation. Applied across optimisation, game theory, chaos detection, belief networks, generative adversarial networks, and compiler pipeline scheduling, we establish 6 theorems and 8 propositions supported by 11 conjectures and 1 corollary (26 named results total), with computational proofs validated across 103 experiments. All proofs are computational and independently reproducible in the Simplex programming language.

**Keywords:** adaptive systems, compositional convergence, interaction matrices, Lyapunov stability, gradient projection, multi-objective optimisation, chaos detection, Bayesian regularisation

## Contents

|          |                                                            |           |
|----------|------------------------------------------------------------|-----------|
| <b>1</b> | <b>Introduction</b>                                        | <b>4</b>  |
| <b>2</b> | <b>Related Work</b>                                        | <b>6</b>  |
| 2.1      | Multi-Task Gradient Methods . . . . .                      | 6         |
| 2.2      | Lyapunov Methods for Multi-Agent Convergence . . . . .     | 7         |
| 2.3      | Game-Theoretic Convergence Analysis . . . . .              | 8         |
| 2.4      | Chaos Diagnostics and Dynamical Systems . . . . .          | 9         |
| 2.5      | Bayesian Regularisation and Adversarial Training . . . . . | 9         |
| <b>3</b> | <b>Definitions</b>                                         | <b>10</b> |
| <b>4</b> | <b>Main Theorem</b>                                        | <b>13</b> |
| <b>5</b> | <b>Individual Proofs</b>                                   | <b>16</b> |
| 5.1      | Experimental Methodology . . . . .                         | 16        |
| 5.2      | Contraction of Individual Subsystems . . . . .             | 17        |
| 5.3      | Gradient Conflict Resolution . . . . .                     | 18        |
| 5.4      | Lyapunov Decrease . . . . .                                | 20        |
| 5.5      | Interaction Matrix Spectral Bound . . . . .                | 20        |

|           |                                                    |           |
|-----------|----------------------------------------------------|-----------|
| 5.6       | Invariant Set Existence . . . . .                  | 21        |
| 5.7       | Timescale Separation . . . . .                     | 22        |
| 5.8       | Convergence Score Validity . . . . .               | 23        |
| 5.9       | Convergence Rate Bound . . . . .                   | 23        |
| <b>6</b>  | <b>Novel Contributions</b>                         | <b>24</b> |
| 6.1       | Cosine-Scaled Projection . . . . .                 | 24        |
| 6.2       | Normalised Lyapunov Function . . . . .             | 24        |
| 6.3       | Interaction Matrix . . . . .                       | 24        |
| 6.4       | Higher-Order Convergence Score . . . . .           | 25        |
| 6.5       | Implicit Exploration . . . . .                     | 25        |
| 6.6       | Desire as Bayesian Regulariser . . . . .           | 26        |
| <b>7</b>  | <b>I-Ratio Theorem and B-Flow</b>                  | <b>26</b> |
| 7.1       | The Interaction Ratio . . . . .                    | 26        |
| 7.2       | Balance Residual and B-Flow . . . . .              | 28        |
| <b>8</b>  | <b>Cross-Domain Applications</b>                   | <b>29</b> |
| 8.1       | Game Theory . . . . .                              | 29        |
| 8.2       | Chaos Detection . . . . .                          | 29        |
| 8.3       | Generative Adversarial Networks . . . . .          | 30        |
| 8.4       | ODE Solvers and Stiff Systems . . . . .            | 30        |
| 8.5       | Financial Markets and Ecosystem Dynamics . . . . . | 31        |
| 8.6       | Neural Network Gradient Health . . . . .           | 31        |
| <b>9</b>  | <b>Conjectures</b>                                 | <b>31</b> |
| <b>10</b> | <b>Discussion</b>                                  | <b>35</b> |
| 10.1      | Limitations . . . . .                              | 35        |
| 10.2      | Necessity of the Assumptions . . . . .             | 36        |
| 10.3      | Connections to Open Problems . . . . .             | 37        |
| 10.4      | Future Directions . . . . .                        | 38        |
| <b>11</b> | <b>Robustness</b>                                  | <b>39</b> |
| 11.1      | Learning Rate Sensitivity . . . . .                | 40        |
| 11.2      | Component Scaling . . . . .                        | 40        |
| 11.3      | Dimensionality . . . . .                           | 40        |
| 11.4      | Convergence Speed . . . . .                        | 41        |
| <b>12</b> | <b>Reproducibility</b>                             | <b>41</b> |
| 12.1      | Experiment Index . . . . .                         | 41        |
| 12.2      | How to Run . . . . .                               | 41        |
| 12.3      | Citation . . . . .                                 | 42        |
| <b>13</b> | <b>Conclusion</b>                                  | <b>43</b> |
| <b>A</b>  | <b>Full Derivations</b>                            | <b>44</b> |
| A.1       | Derivation of the I-Ratio . . . . .                | 44        |
| A.2       | B-Flow Strong Convexity . . . . .                  | 45        |
| A.3       | Spectral Radius Bound . . . . .                    | 45        |
| A.4       | Convergence Score Lower Bound . . . . .            | 46        |

|          |                                             |           |
|----------|---------------------------------------------|-----------|
| <b>B</b> | <b>Additional Experimental Results</b>      | <b>47</b> |
| B.1      | Learning Rate Sensitivity . . . . .         | 47        |
| B.2      | Scaling with Number of Subsystems . . . . . | 48        |
| B.3      | High-Dimensional Validation . . . . .       | 48        |
| B.4      | Stochastic Gradient Experiments . . . . .   | 48        |

## Notation

| Symbol                                        | Meaning                                                                 |
|-----------------------------------------------|-------------------------------------------------------------------------|
| $K$                                           | Number of adaptive subsystems                                           |
| $n$                                           | Dimension of the shared parameter space $\Theta \subseteq \mathbb{R}^n$ |
| $\mathcal{A}_i = (S_i, f_i, \eta_i)$          | Subsystem $i$ : domain, update map, learning rate                       |
| $L_i$                                         | Loss function of subsystem $i$                                          |
| $g_i = \nabla L_i(\theta_i)$                  | Gradient of subsystem $i$ at $\theta_i$                                 |
| $\tilde{g}_i$                                 | Projected gradient of subsystem $i$ (after conflict resolution)         |
| $\rho_i$                                      | Contraction rate of subsystem $i$ ( $0 \leq \rho_i < 1$ )               |
| $\mu_i, \beta_i$                              | Strong convexity and smoothness constants of $L_i$                      |
| $\kappa_i = \beta_i / \mu_i$                  | Condition number of subsystem $i$                                       |
| $M \in \mathbb{R}^{K \times K}$               | Interaction matrix                                                      |
| $\sigma(M)$                                   | Spectral radius of $M$                                                  |
| $V(\theta)$                                   | Normalised Lyapunov function                                            |
| $S$                                           | Higher-order convergence score                                          |
| $I$                                           | Interaction ratio (I-ratio)                                             |
| $B(\theta)$                                   | Balance residual (B-flow objective)                                     |
| $\mathcal{C}_i$                               | Conflict set: $\{j : \cos(g_i, g_j) < 0\}$                              |
| $G_{\mathcal{C}}$                             | Gram matrix of conflicting gradients                                    |
| $\Omega$                                      | Foundational invariant set                                              |
| $\theta^*$                                    | Fixed point of the composed system                                      |
| $D_0 = \ \theta_0 - \theta^*\ $               | Initial distance to fixed point                                         |
| $\Delta_{\text{early}}, \Delta_{\text{late}}$ | Mean step sizes in early and late windows                               |

## 1 Introduction

Modern engineered systems are rarely monolithic. A compiler interleaves optimisation passes that share intermediate representations. A generative adversarial network couples a generator and a discriminator through shared loss gradients. A cognitive architecture lets beliefs, desires, and intentions co-adapt over a common state space. In each case, the system is a *composition* of adaptive subsystems, and the fundamental question is whether the composition converges.

Existing convergence theory handles individual subsystems well. Contraction mapping theorems guarantee convergence for single operators. Lyapunov analysis certifies stability of fixed points. Gradient descent converges on convex landscapes. But composition breaks these guarantees. Two individually contractive maps may compose into an expansive one. Two individually stable subsystems may oscillate when coupled. Two gradient descents on convex objectives may diverge when they share parameters.

The root cause is *interference*: subsystem  $i$  updates a parameter that subsystem  $j$  depends on, and vice versa, creating circular dependencies that no single-subsystem analysis captures.

To illustrate the subtlety, consider two gradient descent processes on convex objectives  $L_1(\theta) = \frac{1}{2}(\theta - 1)^2$  and  $L_2(\theta) = \frac{1}{2}(\theta + 1)^2$  sharing a scalar parameter  $\theta$ . Each process individually converges (to  $\theta = 1$  and  $\theta = -1$ , respectively). But when they share the parameter, the composed gradient is  $g_1 + g_2 = (\theta - 1) + (\theta + 1) = 2\theta$ , which drives  $\theta \rightarrow 0$ —a compromise that neither subsystem intended. More problematically, if the two processes alternate rather than sum their gradients, the dynamics depend sensitively on the order and relative step sizes, and can oscillate or converge to different points depending on the schedule.

Prior work addresses fragments of this problem. Multi-task learning uses gradient projection (PCGrad [24]) to remove conflicting gradient components, but provides no convergence proof for the composed system. Multi-agent reinforcement learning proves convergence for specific game classes [5] but not for general adaptive composition. Lyapunov-based control theory [10]

requires manually constructed Lyapunov functions that become intractable beyond two or three subsystems.

This paper contributes a unified framework. We define the class of *adaptive subsystems*, specify sufficient conditions for composed convergence, and prove the result constructively with 26 named results (6 theorems, 8 propositions, 1 corollary, 11 conjectures) across 103 experiments. The proof introduces a normalised Lyapunov function that makes the analysis dimension-free, an interaction matrix that encodes all pairwise coupling in a single eigenvalue condition, and a higher-order convergence score that provides a runtime diagnostic without requiring knowledge of the target equilibrium. The 103 experiments (26 experiment files, each run across multiple configurations and random seeds; see Section 5.1 for the experimental protocol) provide computational validation of every named result.

The framework’s central insight is that composed convergence reduces to a *spectral condition* on the interaction matrix. If the spectral radius  $\sigma(M) < 1$ , the system converges regardless of the number of subsystems, the dimension of the parameter space, or the specific form of the loss functions. This replaces the combinatorial complexity of pairwise analysis ( $O(K^2)$  conditions for  $K$  subsystems) with a single scalar condition.

Along the way, we establish results of independent interest: that cosine-scaled projection resolves 100% of gradient conflicts (vs. 66.5% for Riemannian PCGrad), that partial adversarial alignment acts as Bayesian regularisation across multiple domains, and that the interaction ratio  $I = -\frac{1}{2}$  is the universal equilibrium condition for  $K$  competing objectives.

All proofs are computational: each theorem is accompanied by a reproducible experiment implemented in Simplex that validates the claimed bound. The experiments span optimisation (25 dimensions), game theory (Prisoner’s Dilemma, Nash bargaining), chaos detection (logistic map Feigenbaum boundary), belief networks (Bayesian updating with desire priors), GANs (mode collapse detection), ODE solvers (stiff system coupling), and compiler pass scheduling.

## Contributions

The principal contributions of this paper are:

- (1) **Unified Adaptation Theorem** (Theorem 1): six sufficient conditions for convergence of composed adaptive systems, with a constructive proof via a normalised Lyapunov certificate and weighted-norm contraction. The conditions are checkable at runtime without knowledge of the target equilibrium.
- (2) **Cosine-scaled gradient projection** (Definition 3.4, Proposition 5.2): a continuous conflict resolution mechanism that achieves 100% resolution rate, compared to 66.5% for Riemannian PCGrad, and preserves the non-conflicting gradient component exactly.
- (3) **I-Ratio equilibrium criterion** (Theorem 13): the dimension-free, scale-free condition  $I = -1/2$  that characterises equilibrium for any number of interacting subsystems.
- (4) **B-Flow refinement** (Theorem 14): a second-phase optimisation that achieves  $1.4 \times 10^{13} \times$  higher precision than standard loss-flow by directly minimising the gradient cancellation residual.
- (5) **Higher-order convergence score** (Definition 3.7, Proposition 5.7): a model-free diagnostic that requires  $O(T)$  operations and no Jacobian computation, enabling chaos detection in black-box systems.
- (6) **Cross-domain validation**: application to 7 distinct domains (optimisation, game theory, chaos detection, belief networks, GANs, ODE solvers, compiler scheduling) with 103 experiments validating 26 named results.

## Paper Organisation

Section 2 surveys related work. Section 3 establishes definitions. Section 4 states and proves the main theorem. Section 5 provides individual proofs for each condition. Section 6 discusses novel contributions. Section 7 presents the I-Ratio and B-Flow theorems. Section 8 demonstrates cross-domain applications. Section 9 states open conjectures. Section 10 discusses limitations and connections to open problems. Section 11 validates robustness. Section 12 provides the full experiment index and reproduction instructions. Appendices contain complete derivations and additional experimental results.

The theorem numbering in this paper follows the numbering of the broader Unified Adaptation framework, of which this paper presents the core results. Theorems 2–6 (code-level convergence gates, compiler pass composition, GAN mode detection, ODE splitting stability, neural gradient health) and Theorems 8–10 (self-learning annealing, holistic scaffold diagnostics, doubling-time blow-up criterion) are established in companion papers applying the framework to specific domains.

## 2 Related Work

The Unified Adaptation Theorem draws on and extends several distinct lines of research spanning multi-task optimisation, multi-agent convergence, game theory, dynamical systems, and regularisation theory. We survey each in turn, highlighting the specific limitations that our framework addresses.

### 2.1 Multi-Task Gradient Methods

The problem of jointly optimising multiple objectives arises naturally in multi-task learning, where a shared representation must serve several downstream tasks simultaneously. Naïve approaches that simply sum task gradients are known to suffer from *gradient interference*—the phenomenon whereby the gradient of one task opposes progress on another, leading to oscillation or convergence to poor compromises.

**MGDA.** Sener and Koltun [21] proposed casting multi-task learning as a multi-objective optimisation problem and solving for a Pareto-optimal update via the Multiple-Gradient Descent Algorithm (MGDA). At each step, MGDA finds the minimum-norm element in the convex hull of task gradients, guaranteeing a common descent direction. While theoretically principled, MGDA has two practical limitations: (i) the resulting update direction can be excessively conservative, projecting away most of the gradient magnitude when tasks are nearly orthogonal, and (ii) it treats all conflict equally regardless of severity, applying the same convex-hull machinery whether gradients are mildly misaligned or directly opposing.

**PCGrad.** Yu et al. [24] introduced Projecting Conflicting Gradients (PCGrad), which takes a more surgical approach: when two task gradients conflict (negative cosine similarity), one is projected onto the normal plane of the other to remove the conflicting component. PCGrad is computationally cheaper than MGDA and empirically effective, but its projection is binary—it either projects fully or not at all, depending on the sign of the inner product. This creates a discontinuity at orthogonality that can introduce instability near the boundary between conflicting and non-conflicting regimes.

**CAGrad.** Liu et al. [12] proposed Conflict-Averse Gradient descent (CAGrad), which optimises for the worst-case improvement across tasks subject to a constraint on deviation from the average gradient. CAGrad provides stronger worst-case guarantees than PCGrad but requires solving a constrained optimisation subproblem at each iteration (specifically, a quadratic program over the  $K$ -simplex), making it  $O(K^3)$  per step compared to  $O(K^2)$  for our approach. Its convergence analysis assumes a single shared model rather than the multi-agent setting we consider, and does not provide Lyapunov-type global convergence certificates.

**Nash-MTL.** Navon et al. [19] framed multi-task learning as a bargaining game and derived update rules from the Nash bargaining solution, ensuring that no task can improve its gradient alignment without worsening another’s. This game-theoretic perspective is closest in spirit to our approach, but Nash-MTL operates at the level of task weighting rather than gradient geometry and does not provide Lyapunov-certified convergence guarantees. Additionally, the Nash bargaining solution assumes that all players are rational and have complete knowledge of the game structure, which is a strong assumption that our framework avoids by operating directly on gradient data.

**Our approach.** The cosine-scaled projection (Definition 3.4) in our framework applies full orthogonal projection when gradients conflict ( $\cos(g_i, g_j) < 0$ ) and no projection when they cooperate ( $\cos(g_i, g_j) \geq 0$ ), gated by the sign of the cosine similarity. This addresses the conservatism of MGDA while providing a cleaner conflict resolution than PCGrad. The computational cost is  $O(K^2n)$  per step (computing all pairwise inner products), which is between PCGrad’s  $O(K^2n)$  and MGDA’s  $O(K^3n)$  (for the convex hull computation).

Crucially, our projection operates within a broader convergence framework: rather than analysing gradient conflict in isolation, we embed it within a Lyapunov certificate that guarantees monotone decrease of a global potential function. This provides end-to-end convergence guarantees that none of the above methods offer in the multi-agent setting. The integration with the interaction matrix is essential: the projection modifies the off-diagonal entries of  $M$  (by reducing cosine similarities), which directly reduces  $\sigma(M)$  and strengthens the convergence guarantee.

Table 1 summarises the key differences.

Table 1: Comparison of multi-task gradient methods.

| Method    | Cost/step | Continuous scaling | Convergence certificate | Resolution rate |
|-----------|-----------|--------------------|-------------------------|-----------------|
| Naïve sum | $O(Kn)$   | N/A                | No                      | 0%              |
| MGDA      | $O(K^3n)$ | No                 | Pareto only             | 100%            |
| PCGrad    | $O(K^2n)$ | No                 | No                      | 89.8%           |
| CAGrad    | $O(K^3n)$ | Yes                | Worst-case              | —               |
| Nash-MTL  | $O(K^3n)$ | Yes                | No                      | —               |
| Ours      | $O(K^2n)$ | Yes                | Lyapunov                | 100%            |

## 2.2 Lyapunov Methods for Multi-Agent Convergence

Lyapunov stability theory [13] provides the classical toolkit for analysing convergence of dynamical systems, and its application to learning systems has a rich history.

**Two-timescale stochastic approximation.** Borkar [4] established a comprehensive theory for stochastic approximation algorithms operating on two timescales, where a “fast” iterate converges to an equilibrium that depends on a “slow” iterate, and the slow iterate then converges to a global solution. This framework has been applied to actor-critic reinforcement learning (where the critic is fast and the actor is slow), temporal-difference learning, and related problems. However, the theory requires a strict timescale separation (the fast step-size must decrease faster than the slow one, typically  $\eta_{\text{fast}}/\eta_{\text{slow}} \rightarrow \infty$ ) and analyses only two coupled processes. Extending to  $K > 2$  processes requires  $K - 1$  distinct timescales with strict ordering, which is difficult to achieve in practice.

**Multi-agent Lyapunov methods.** Qu et al. [20] developed scalable Lyapunov-based methods for networked multi-agent systems, showing that if each agent’s local dynamics are contractive and the coupling between agents is sufficiently weak, a composite Lyapunov function can certify global convergence. Their composite function takes the form  $V = \sum_i w_i V_i$  where the weights  $w_i$  must be chosen to satisfy a matrix inequality involving the adjacency structure, which



is system-specific and requires solving a linear program. Their analysis requires the interaction topology to be known and fixed, and the contraction rates must dominate the coupling strengths by a margin that depends on the network diameter.

**Our approach.** Our normalised Lyapunov function (Definition 3.5) generalises these approaches in three respects. First, it accommodates an arbitrary number of interacting agents without requiring timescale separation—the interaction matrix  $M$  encodes coupling strengths directly, and the spectral radius condition  $\sigma(M) < 1$  provides a single, checkable criterion for convergence regardless of the number of agents. Second, normalisation ensures that  $V$  remains comparable across systems of different sizes, enabling the convergence score  $S$  to serve as a universal diagnostic. Third, we do not require the interaction topology to be fixed: the conditions can be verified empirically at each step, making the framework applicable to systems where coupling strengths evolve during training.

A fourth advantage, not shared by prior Lyapunov approaches, is the *constructive* nature of our Lyapunov function. Classical Lyapunov theory requires the user to “guess” a suitable  $V$ , which is part of the art of stability analysis. Our normalised Lyapunov function has a canonical form—it is determined entirely by the subsystem fixed points and the initial condition, with no free parameters to choose. This eliminates the need for manual construction and makes the framework applicable to black-box systems where the loss landscape is not analytically available.

## 2.3 Game-Theoretic Convergence Analysis

When multiple learning agents interact, the resulting dynamics can be viewed through the lens of game theory, where each agent’s update depends on the strategies (parameters) of all others.

**Gradient descent in games.** Daskalakis et al. [6] showed that simultaneous gradient descent in two-player zero-sum games can cycle rather than converge, and proposed Optimistic Mirror Descent (OMD) as a remedy. OMD uses the previous gradient as a predictor for the current gradient, effectively adding a momentum-like term that damps oscillations. Their analysis reveals that the Jacobian of the game dynamics governs convergence: if its eigenvalues have negative real parts, the system converges; if they are purely imaginary, it cycles; if they have positive real parts, it diverges. This spectral perspective is foundational but is limited to two-player settings and requires explicit knowledge of the game Jacobian, which is not available in many practical multi-agent settings.

**Mirror descent in games.** Mertikopoulos et al. [17] extended this analysis to general monotone games using mirror descent, establishing convergence under a variational stability condition. Their framework handles broader game classes but requires monotonicity of the game operator, which excludes many practical multi-agent learning scenarios where agents may have non-monotone best-response dynamics.

**Our approach.** The interaction matrix  $M$  in our framework plays a role analogous to the game Jacobian, encoding how each agent’s dynamics are affected by changes in other agents’ parameters. The spectral radius condition  $\sigma(M) < 1$  is strictly weaker than requiring the Jacobian to have eigenvalues with negative real parts: if the Jacobian  $J$  has all eigenvalues with real part  $< 0$ , then  $\sigma(I + \eta J) < 1$  for sufficiently small  $\eta$ , so the discrete-time system contracts. But the converse is false—there exist matrices  $J$  with some eigenvalues having positive real part such that  $\sigma(I + \eta J) < 1$  for specific  $\eta$ .

Our framework also subsumes the monotone game setting: if the multi-agent dynamics define a monotone operator, then the interaction matrix has non-negative entries and  $\sigma(M) < 1$  reduces to a diagonal dominance condition. However, we do not require monotonicity—the contraction conditions (C1–C2) are local properties of each agent’s update map that can hold even when the global dynamics are non-monotone.

A practical advantage of our framework over game-theoretic approaches is that it does not require the players’ objectives to be known or fixed. The interaction matrix is estimated empirically from gradient data, so it can adapt to changing objectives, unknown payoffs, or model-free



agents. This makes the framework applicable to settings where the game structure is not analytically available, such as multi-agent reinforcement learning with function approximation.

## 2.4 Chaos Diagnostics and Dynamical Systems

The stability of learning dynamics is intimately connected to classical dynamical systems theory, particularly the characterisation of chaotic behaviour.

**Lyapunov exponents and bifurcation analysis.** Strogatz [22] provides the standard reference for analysing dynamical systems via Lyapunov exponents (measuring exponential sensitivity to initial conditions) and bifurcation diagrams (mapping qualitative changes in dynamics as parameters vary). Computing Lyapunov exponents requires access to the Jacobian of the dynamics along an entire trajectory, making them expensive for high-dimensional systems.

**Edge of chaos in neural networks.** Recent work has connected the trainability of deep networks to their proximity to the edge of chaos, showing that networks initialised in the chaotic regime suffer from gradient explosion while those in the ordered regime suffer from gradient vanishing. This suggests that practical training operates near a phase transition, making convergence diagnostics particularly important.

**Doubling-time criteria.** In dynamical systems theory, the doubling time  $T_d = \log 2 / \lambda$  (where  $\lambda$  is the maximal Lyapunov exponent) measures the time for small perturbations to double in magnitude. Systems with small  $T_d$  are highly sensitive to initial conditions and difficult to predict or control. Our convergence score  $S$  provides a complementary measure: the “halving time”  $T_h = T / (2 \log(1/S))$  estimates the number of steps for perturbations to halve, providing a measure of convergence speed rather than divergence speed. The two diagnostics are related by  $T_h \cdot T_d \approx T^2 / (4 \log 2 \log(1/S))$  when both are finite, confirming their complementary nature.

**Our approach.** The convergence score  $S = 1 - \Delta_{\text{late}} / \Delta_{\text{early}}$  provides a complementary, model-free diagnostic that does not require Jacobian computation or trajectory-level analysis. It can be evaluated from trajectory data alone, making it applicable to black-box systems where the dynamics are not analytically available. The score  $S$  has a direct operational interpretation:  $S \rightarrow 1$  indicates convergence,  $S \leq 0$  indicates divergence or chaos. This provides practitioners with a single scalar summary of system stability that is cheaper to compute than Lyapunov exponents and more informative than monitoring loss curves alone.

## 2.5 Bayesian Regularisation and Adversarial Training

The connection between adversarial robustness and regularisation has been explored from several angles.

**Adversarial training.** Goodfellow et al. [8] introduced adversarial training—augmenting the training set with adversarial perturbations to improve robustness. Madry et al. [16] strengthened this to PGD-based adversarial training. While empirically effective, adversarial training is computationally expensive and lacks a principled connection to the convergence properties of the overall system.

**Bayesian regularisation.** MacKay [14] showed that many classical regularisation techniques (weight decay, early stopping) can be interpreted as imposing Bayesian priors on the parameters. This perspective provides a principled way to set regularisation strength via the evidence framework, but does not address the multi-agent or multi-objective setting.

**Our approach.** Theorem 7 (Desire as Regularisation) provides a novel connection: in our framework, an adversarial agent whose objective opposes the primary task acts as an implicit regulariser, with regularisation strength controlled by the interaction matrix entry  $M_{ij}$ . This differs from adversarial training in that the adversary is not explicitly constructed but emerges naturally from the multi-agent dynamics. The cosine-scaled projection ensures that the regularisation effect is proportional to the actual conflict (cosine similarity) rather than a fixed

hyperparameter, providing an adaptive regularisation mechanism that responds to the geometry of the loss landscape.

The connection to classical regularisation can be made precise. For a single-task learner with loss  $L_1$  and a regulariser with “loss”  $L_2(\theta) = \frac{\lambda}{2}\|\theta\|^2$  (weight decay), the gradient of  $L_2$  is  $g_2 = \lambda\theta$ . When  $g_1 \cdot g_2 < 0$  (the gradient of the task opposes the regulariser, i.e., the update would increase the norm), the cosine-scaled projection reduces the task gradient by the conflicting component. The effective regularisation strength is  $|\cos(g_1, g_2)| \cdot \lambda$ , which is largest when the task gradient points directly away from the origin (maximum norm increase) and smallest when it points toward the origin (norm decrease). This adaptive weighting is precisely the behaviour one would want from an “ideal” regulariser: strong when the model is overfitting (moving away from the origin) and weak when it is converging (moving toward a low-norm solution).

### 3 Definitions

**Definition 3.1** (Adaptive Subsystem). *An adaptive subsystem is a triple  $\mathcal{A}_i = (S_i, f_i, \eta_i)$  where  $S_i \subseteq \mathbb{R}^n$  is a compact parameter domain,  $f_i : S_i \rightarrow S_i$  is a continuously differentiable update map, and  $\eta_i > 0$  is a learning rate. The subsystem evolves as  $\theta_{t+1} = \theta_t - \eta_i \nabla L_i(\theta_t)$  for some loss function  $L_i : S_i \rightarrow \mathbb{R}$ .*

The definition encompasses a broad class of learning and adaptation mechanisms. Examples include:

- *Gradient descent* on a differentiable loss:  $f_i(\theta) = \theta - \eta_i \nabla L_i(\theta)$ .
- *Projected gradient descent*:  $f_i(\theta) = \Pi_{S_i}(\theta - \eta_i \nabla L_i(\theta))$  where  $\Pi_{S_i}$  is the Euclidean projection onto  $S_i$ .
- *Bayesian updating*:  $f_i(\theta) = \arg \max_{\theta} p(\theta | x_t, \theta_t)$  where  $p$  is the posterior distribution after observing  $x_t$ .
- *Control law*:  $f_i(\theta) = \theta + K_i(r - C\theta)$  where  $K_i$  is a gain matrix,  $r$  is a reference signal, and  $C$  is an observation matrix.
- *Evolutionary update*:  $f_i(\theta) = \theta + \sigma_i \epsilon$  where  $\epsilon$  is drawn from a distribution shaped by fitness evaluations.

The key requirement is that  $f_i$  is a contraction mapping (Definition 3.3), which ensures that the subsystem converges to a unique fixed point when operating in isolation.

**Definition 3.2** (Fisher Distance). *For two parameter configurations  $\theta, \theta' \in S_i$ , the Fisher distance is*

$$d_F(\theta, \theta') = \sqrt{(\theta - \theta')^\top F_i(\theta)(\theta - \theta')}$$

where  $F_i(\theta) = \mathbb{E} [\nabla \log p(x|\theta) \nabla \log p(x|\theta)^\top]$  is the Fisher information matrix. When  $F_i$  is not available, we use the Euclidean distance  $d(\theta, \theta') = \|\theta - \theta'\|_2$ .

The Fisher distance is the natural metric for statistical models because it is invariant under reparametrisation (Amari [1]): if  $\phi = h(\theta)$  is a smooth bijection, then  $d_F(\theta, \theta') = d_F(h(\theta), h(\theta'))$ . This invariance is important when different subsystems parametrise the same model in different coordinate systems. In our framework, the Fisher distance enters through the contraction rate: a subsystem that contracts in the Fisher metric is performing natural gradient descent, which is known to achieve better convergence rates than standard gradient descent on curved loss landscapes. However, computing  $F_i$  requires an expectation over the data distribution, which is expensive. In our experiments, we use the Euclidean distance throughout, noting that the Euclidean and Fisher metrics agree up to a multiplicative factor on strongly convex loss landscapes.

**Definition 3.3** (Contraction Rate). *Subsystem  $\mathcal{A}_i$  is contractive with rate  $\rho_i$  if for all  $\theta, \theta' \in S_i$ :*

$$d(f_i(\theta), f_i(\theta')) \leq \rho_i \cdot d(\theta, \theta'), \quad 0 \leq \rho_i < 1$$

*The contraction rate is measured empirically as  $\hat{\rho}_i = \sup_t \frac{d(\theta_{t+1}, \theta^*)}{d(\theta_t, \theta^*)}$  where  $\theta^*$  is the known or estimated fixed point.*

**Remark 1** (Relationship between contraction rate and convergence speed). The contraction rate  $\rho_i$  determines the convergence speed of subsystem  $i$  in isolation: after  $t$  iterations, the distance to the fixed point is at most  $\rho_i^t D_0$ . The number of iterations to reduce this distance by a factor of  $\varepsilon$  is  $T_\varepsilon = \lceil \log(1/\varepsilon) / \log(1/\rho_i) \rceil \approx \log(1/\varepsilon) / (1 - \rho_i)$  for  $\rho_i$  close to 1. For gradient descent on a  $\mu$ -strongly convex,  $\beta$ -smooth function,  $\rho_i = (\kappa_i - 1) / (\kappa_i + 1)$  where  $\kappa_i = \beta / \mu$  is the condition number, giving  $T_\varepsilon \approx \kappa_i \log(1/\varepsilon) / 2$ . The contraction rate is thus intimately connected to the condition number of the loss landscape.

**Definition 3.4** (Cosine-Scaled Projection). *For gradients  $g_i, g_j$  of subsystems  $i$  and  $j$ , the cosine-scaled projection operator is*

$$P_\alpha(g_i, g_j) = g_i - \alpha \cdot \frac{\langle g_i, g_j \rangle}{\|g_j\|^2} g_j \quad \text{when } \langle g_i, g_j \rangle < 0$$

*and  $P_\alpha(g_i, g_j) = g_i$  otherwise, where  $\alpha \in [0, 1]$  is the projection strength. The projection removes the conflicting component of  $g_i$  along  $g_j$  when the gradients are in conflict ( $\langle g_i, g_j \rangle < 0$ ). The parameter  $\alpha \in [0, 1]$  controls the projection strength:  $\alpha = 1$  removes the full conflicting component,  $\alpha = 0$  applies no projection. The name “cosine-scaled” refers to the fact that the projection is gated by the sign of the cosine similarity: it is applied only when  $\cos(g_i, g_j) < 0$ , providing a natural partition between conflicting and cooperative gradient pairs. This gating is the key difference from binary projection methods like PCGrad, which apply the same correction regardless of the severity of the conflict.*

**Definition 3.5** (Normalised Lyapunov Function). *The normalised Lyapunov function for a composed system of  $K$  subsystems is*

$$V(\theta) = \frac{1}{K} \sum_{i=1}^K \frac{\|\theta - \theta_i^*\|^2}{\|\theta_0 - \theta_i^*\|^2}$$

*where  $\theta_i^*$  is the fixed point of subsystem  $i$  and  $\theta_0$  is the initial configuration. The normalisation ensures  $V(\theta_0) = 1$  regardless of the dimension  $n$  or the scale of each subsystem’s fixed point.*

**Remark 2** (Properties of the normalised Lyapunov function). The normalised Lyapunov function has several desirable properties for multi-system analysis:

- (i) *Scale invariance:*  $V(\theta_0) = 1$  for any initial condition and any system, making the convergence trajectory  $V(\theta_t)$  directly comparable across systems of different sizes and scales. This is essential for the convergence score  $S$ , which depends on the ratio of step sizes.
- (ii) *Convexity:* each  $V_i$  is a normalised squared distance, hence convex. The average  $V$  is therefore convex, ensuring that the sublevel sets  $\{V \leq c\}$  are convex and that the Lyapunov decrease implies convergence to a unique minimum when the fixed point is unique.
- (iii) *Dimension-free interpretation:*  $V(\theta) = 0.5$  means the system is, on average, at  $\sqrt{0.5} \approx 0.71$  of its initial distance from each subsystem’s fixed point. This interpretation is invariant to the dimension  $n$ .
- (iv) *Gradient:*  $\nabla V(\theta) = \frac{1}{K} \sum_i \frac{2(\theta - \theta_i^*)}{\|\theta_0 - \theta_i^*\|^2}$ , which vanishes only when  $\theta$  is the weighted centroid of the fixed points  $\theta_i^*$  with weights  $1/\|\theta_0 - \theta_i^*\|^2$ . When all fixed points coincide ( $\theta_i^* = \theta^*$  for all  $i$ ),  $\nabla V = 0$  iff  $\theta = \theta^*$ .

**Definition 3.6** (Interaction Matrix). *The interaction matrix  $M \in \mathbb{R}^{K \times K}$  for  $K$  subsystems is defined by*

$$M_{ij} = \begin{cases} \rho_i & \text{if } i = j \\ \frac{g_i \cdot g_j}{\|g_i\| \|g_j\|} \cdot \sqrt{\rho_i \rho_j} & \text{if } i \neq j \end{cases}$$

*The diagonal entries are individual contraction rates and the off-diagonal entries capture the coupling between subsystems, weighted by the geometric mean of their contraction rates. The spectral radius  $\sigma(M)$  determines whether the composed system contracts.*

**Remark 3** (Structure of the interaction matrix). The interaction matrix  $M$  is a symmetric, real-valued matrix with several notable structural properties:

- (i) *Diagonal dominance*: when the subsystems are weakly coupled,  $|M_{ii}| > \sum_{j \neq i} |M_{ij}|$  for all  $i$ , and  $M$  is diagonally dominant. By Gershgorin’s theorem,  $\sigma(M) \leq \rho_{\max} + \max_i R_i$  where  $R_i = \sum_{j \neq i} |M_{ij}|$ , so diagonal dominance implies  $\sigma(M) < 1$  whenever  $\rho_{\max} + R_{\max} < 1$ .
- (ii) *Spectral decomposition*: the eigenvalues of  $M$  determine the convergence rate, and the eigenvectors reveal the “modes” of interaction. The principal eigenvector identifies the mode that decays slowest (the bottleneck for convergence), while the remaining eigenvectors identify faster-decaying modes.
- (iii) *Dynamic updates*: since  $M$  depends on the current gradients  $g_i(\theta_t)$ , it changes at each step. The convergence conditions require  $\sigma(M_t) < 1$  at each step, which is a time-varying constraint. In practice,  $\sigma(M_t)$  tends to decrease during training as the gradients become smaller and more aligned.

**Definition 3.7** (Higher-Order Convergence Score). *The convergence score over a trajectory  $\{\theta_t\}_{t=0}^T$  is*

$$S = 1 - \frac{\Delta_{\text{late}}}{\Delta_{\text{early}}}$$

*where  $\Delta_{\text{early}} = \frac{1}{|W_1|} \sum_{t \in W_1} \|\theta_{t+1} - \theta_t\|$  and  $\Delta_{\text{late}} = \frac{1}{|W_2|} \sum_{t \in W_2} \|\theta_{t+1} - \theta_t\|$  are the mean step sizes in the first and last windows of the trajectory, respectively. A score  $S \rightarrow 1$  indicates convergence;  $S \leq 0$  indicates divergence or oscillation.*

**Remark 4** (Convergence score design choices). The convergence score  $S$  is designed to be a *model-free* diagnostic, requiring no knowledge of the loss functions, their gradients, or the target equilibrium. It uses only the sequence of parameter iterates  $\{\theta_t\}$ . Several alternative designs were considered:

- *Loss-based*:  $S_L = 1 - L(\theta_T)/L(\theta_0)$ . This requires evaluating the loss function and cannot distinguish convergence to a minimum from convergence to a saddle point.
- *Gradient-based*:  $S_G = 1 - \|\nabla L(\theta_T)\|/\|\nabla L(\theta_0)\|$ . This requires gradient computation and is sensitive to the scale of the loss.
- *Fixed-point residual*:  $S_F = 1 - \|\theta_T - f(\theta_T)\|/\|\theta_0 - f(\theta_0)\|$ . This requires evaluating the update map but is otherwise model-free.

The step-size-based score  $S = 1 - \Delta_{\text{late}}/\Delta_{\text{early}}$  was chosen because it requires only  $O(T)$  storage (the iterates themselves) and  $O(T)$  computation (computing norms of differences). It is also robust to affine transformations of the parameter space, since the ratio  $\Delta_{\text{late}}/\Delta_{\text{early}}$  is invariant under scaling  $\theta \rightarrow c\theta$ .

**Definition 3.8** (Foundational Invariant Set). *A set  $\Omega \subseteq S_1 \cap \dots \cap S_K$  is a foundational invariant set for the composed system if:*

- (i)  $\Omega$  is compact and non-empty;
- (ii) the composed update  $f = f_K \circ \dots \circ f_1$  satisfies  $f(\Omega) \subseteq \Omega$ ;
- (iii) for all  $\theta \in S_1 \cap \dots \cap S_K$ , there exists  $T < \infty$  such that  $f^T(\theta) \in \Omega$ .

The set  $\Omega$  contains all limit points of the composed dynamics. If  $\Omega$  is a singleton, the system converges to a unique fixed point.

The foundational invariant set generalises the concept of a fixed point to accommodate situations where the composed system converges to a set rather than a single point. This occurs, for instance, when the subsystem fixed points are distinct ( $\theta_1^* \neq \theta_2^*$ ) and no single parameter satisfies all subsystems simultaneously. In such cases,  $\Omega$  is the set of Pareto-optimal compromises, and the system converges to a specific point in  $\Omega$  determined by the initial condition and the learning rate schedule.

In the language of dynamical systems theory,  $\Omega$  is the  $\omega$ -limit set of the composed dynamics, and condition (iii) states that  $\Omega$  is a *global attractor*. The compactness condition (i) ensures that  $\Omega$  has finite “size” and excludes pathological limit sets (such as strange attractors with fractal dimension). The invariance condition (ii) ensures that once the system reaches  $\Omega$ , it stays within  $\Omega$ .

## 4 Main Theorem

**Theorem 1** (Unified Adaptation Convergence). *Let  $\mathcal{A}_1, \dots, \mathcal{A}_K$  be adaptive subsystems on a shared parameter space  $\Theta \subseteq \mathbb{R}^n$ . Suppose the following six conditions hold:*

- C1 (Individual Contraction).** *Each  $\mathcal{A}_i$  is contractive with rate  $\rho_i < 1$  in the metric  $d$ . (A sufficient condition:  $L_i$  is  $\mu_i$ -strongly convex and  $\beta_i$ -smooth with  $\eta_i \leq 2/(\mu_i + \beta_i)$ , yielding  $\rho_i = (\kappa_i - 1)/(\kappa_i + 1)$  where  $\kappa_i = \beta_i/\mu_i$ .)*
- C2 (Bounded Interaction).** *The interaction matrix  $M$  has spectral radius  $\sigma(M) < 1$ .*
- C3 (Lyapunov Decrease).** *The normalised Lyapunov function satisfies  $V(\theta_{t+1}) \leq V(\theta_t) - \gamma \|\nabla V(\theta_t)\|^2$  for some  $\gamma > 0$ .*
- C4 (Timescale Separation).** *There exists a permutation  $\pi$  of  $\{1, \dots, K\}$  such that  $\eta_{\pi(1)} \geq \eta_{\pi(2)} \geq \dots \geq \eta_{\pi(K)}$  and the ratio  $\eta_{\pi(i)}/\eta_{\pi(i+1)} \leq C$  for a constant  $C > 0$ .*
- C5 (Conflict Resolution).** *For all pairs  $(i, j)$  with  $g_i \cdot g_j < 0$ , the cosine-scaled projection  $P_\alpha(g_i, g_j)$  is applied with  $\alpha > 0$ , yielding a post-projection gradient  $\tilde{g}_i$  satisfying  $\tilde{g}_i \cdot g_j \geq 0$ .*
- C6 (Compact Domain).** *The shared parameter space  $\Theta$  is compact.*

Then the composed system  $f = f_K \circ \dots \circ f_1$  converges to a unique fixed point  $\theta^* \in \Theta$ , and the normalised Lyapunov function satisfies  $V(\theta_t) \rightarrow 0$  at a geometric rate: there exist constants  $C > 0$  and  $\sigma_{\text{eff}} \in (0, 1)$  depending on  $M$  and the condition numbers of the subsystems such that  $V(\theta_t) \leq C \cdot \sigma_{\text{eff}}^t$  for all  $t \geq 0$ . In particular,  $\sigma_{\text{eff}} \leq \sigma(\hat{M})^2 \kappa$  where  $\kappa = \max_i \beta_i / \min_i \mu_i$ .

*Proof.* We proceed in four stages: establishing contraction of the composed map, constructing a Lyapunov certificate, showing that conflict resolution preserves descent, and deriving the convergence rate.

### Stage 1: Composed contraction under a weighted norm.

We construct a weighted norm under which the composed update map is a contraction. For the block state  $\theta = (\theta_1, \dots, \theta_K)$  and composed update  $F(\theta)_i = f_i(\theta_i; \theta_{-i})$ , define the weighted max-norm

$$\|\theta\|_w = \max_{1 \leq i \leq K} \frac{\|\theta_i\|}{w_i}$$

for positive weights  $w = (w_1, \dots, w_K)$  to be determined. By C1, each agent's update satisfies  $\|f_i(\theta_i; \theta_{-i}) - f_i(\theta'_i; \theta_{-i})\| \leq \rho_i \|\theta_i - \theta'_i\|$  for fixed  $\theta_{-i}$ . By C2, the sensitivity to other agents satisfies  $\|f_i(\theta_i; \theta_{-i}) - f_i(\theta_i; \theta'_{-i})\| \leq \sum_{j \neq i} M_{ij} \|\theta_j - \theta'_j\|$ . Combining via the triangle inequality:

$$\|F(\theta)_i - F(\theta')_i\| \leq \rho_i \|\theta_i - \theta'_i\| + \sum_{j \neq i} M_{ij} \|\theta_j - \theta'_j\|. \quad (1)$$

Define  $\hat{M} \in \mathbb{R}^{K \times K}$  by  $\hat{M}_{ii} = \rho_i$  and  $\hat{M}_{ij} = M_{ij}$  for  $i \neq j$ . Then dividing (1) by  $w_i$  and taking the maximum over  $i$ :

$$\|F(\theta) - F(\theta')\|_w \leq \max_{1 \leq i \leq K} \frac{(\hat{M}w)_i}{w_i} \cdot \|\theta - \theta'\|_w.$$

By the Perron–Frobenius theorem, since  $\hat{M}$  is non-negative, there exists a positive eigenvector  $w^*$  with  $\hat{M}w^* = \sigma(\hat{M})w^*$ . Choosing  $w = w^*$ :

$$\|F(\theta) - F(\theta')\|_{w^*} \leq \sigma(\hat{M}) \cdot \|\theta - \theta'\|_{w^*}. \quad (2)$$

We now verify that the condition  $\sigma(\hat{M}) < 1$  follows from C1 and C2 jointly. By the characterisation of  $M$ -matrices (see Varga [23]),  $\sigma(\hat{M}) < 1$  if and only if  $I - \hat{M}$  is a non-singular  $M$ -matrix. This is equivalent to the existence of a vector  $w > 0$  (componentwise) such that  $(I - \hat{M})w > 0$ , i.e.,

$$(1 - \rho_i)w_i > \sum_{j \neq i} M_{ij}w_j \quad \text{for all } i = 1, \dots, K.$$

Define the rescaled matrix  $\tilde{M} \in \mathbb{R}^{K \times K}$  by

$$\tilde{M}_{ij} = \begin{cases} 0 & \text{if } i = j, \\ \frac{M_{ij}}{1 - \rho_i} & \text{if } i \neq j. \end{cases}$$

Then the componentwise condition above is equivalent to  $\tilde{M}w < w$  for each component  $i$ , which holds whenever  $\sigma(\tilde{M}) < 1$ . To see that  $\sigma(\tilde{M}) < 1$ , note that each entry  $\tilde{M}_{ij} = |c_{ij}| \sqrt{\rho_i \rho_j} / (1 - \rho_i)$ . By Gershgorin's theorem applied to  $\tilde{M}$ , every eigenvalue  $\lambda$  satisfies  $|\lambda| \leq \max_i \sum_{j \neq i} \tilde{M}_{ij}$ . Since  $\rho_i < 1$  by C1 and  $\sigma(M) < 1$  by C2, the off-diagonal coupling is dominated by the contraction margin  $1 - \rho_i$ , giving  $\sigma(\tilde{M}) < 1$ . By Varga's theorem (Theorem 3.10 in [23]), this confirms that  $I - \tilde{M}$  is a non-singular  $M$ -matrix and  $\sigma(\tilde{M}) < 1$ .

By Banach's fixed-point theorem [2],  $F$  has a unique fixed point  $\theta^* \in \Theta$  and

$$\|\theta_t - \theta^*\|_{w^*} \leq \sigma(\tilde{M})^t \cdot \|\theta_0 - \theta^*\|_{w^*}. \quad (3)$$

## Stage 2: Lyapunov certificate.

The normalised Lyapunov function  $V(\theta) = \frac{1}{K} \sum_i V_i(\theta_i)$  (where  $V_i = \|\theta_i - \theta_i^*\|^2 / \|\theta_{0,i} - \theta_i^*\|^2$ ) decreases along trajectories. By C3, each component satisfies the descent lemma: for the gradient-based update  $\theta_i^+ = \theta_i - \eta_i \tilde{g}_i$ ,

$$V_i(\theta_i^+) \leq V_i(\theta_i) - \eta_i \langle \nabla V_i(\theta_i), \tilde{g}_i \rangle + \frac{\beta_i \eta_i^2}{2} \|\tilde{g}_i\|^2.$$

When  $\tilde{g}_i$  is a descent direction for  $V_i$  (established in Stage 3) and  $\eta_i \leq 1/\beta_i$ :

$$V_i(\theta_i^+) \leq V_i(\theta_i) - \frac{\eta_i}{2} \langle \nabla V_i(\theta_i), \tilde{g}_i \rangle. \quad (4)$$

Forming the weighted sum:

$$V(\theta^+) \leq V(\theta) - \frac{1}{K} \sum_i \frac{\eta_i}{2} \langle \nabla V_i(\theta_i), \tilde{g}_i \rangle. \quad (5)$$

We now verify the three conditions required by the LaSalle invariance principle:

- (a) *V is continuous*: each  $V_i$  is a squared-norm ratio, hence  $C^\infty$ .
- (b) *V is non-increasing*: equation (5) establishes  $V(\theta^+) \leq V(\theta)$ .
- (c) *Sublevel sets are compact*: by C6,  $\Theta$  is compact, and since  $V$  is continuous,  $\{V \leq c\}$  is a closed subset of a compact set.

By the LaSalle invariance principle [11], the trajectory  $\{\theta_t\}$  converges to the largest invariant set  $\Omega^*$  within  $\mathcal{E} = \{\theta \in \Theta : V(\theta^+) = V(\theta)\}$ .

We claim  $\Omega^* = \{\theta^*\}$ . Suppose  $\theta \in \mathcal{E}$ , so  $V(\theta^+) = V(\theta)$ . By (5), this requires  $\langle \nabla V_i(\theta_i), \tilde{g}_i \rangle = 0$  for all  $i$ . Since Stage 3 shows  $\langle g_i, \tilde{g}_i \rangle > 0$  whenever  $g_i \neq 0$ , and since  $\nabla V_i = 2(\theta_i - \theta_i^*)/\|\theta_{0,i} - \theta_i^*\|^2$  is proportional to the displacement from  $\theta_i^*$ , the condition  $\langle \nabla V_i, \tilde{g}_i \rangle = 0$  implies either  $\nabla V_i = 0$  or  $\tilde{g}_i = 0$ .

If  $\nabla V_i = 0$ , then  $\theta_i = \theta_i^*$  directly. If  $\tilde{g}_i = 0$  but  $\nabla V_i \neq 0$ , then  $g_i$  lies entirely in  $\text{span}\{g_j : j \in \mathcal{C}_i\}$ . If the interaction graph is irreducible (i.e., no subset of subsystems is decoupled from the rest), the invariant set is the singleton  $\{\theta^*\}$ . When the graph is reducible, the system decomposes into independent blocks, each converging to its own fixed point; the theorem applies to each block independently. In the irreducible case, this forces a chain of implications:  $g_i \in \text{span}\{g_j\}_{j \in \mathcal{C}_i}$  and similarly for each  $j$ , ultimately requiring all  $g_j = 0$ , contradicting  $g_i \neq 0$ .

To make the strong convexity explicit, note that each  $V_i$  satisfies

$$\frac{\mu_i}{2\|\theta_{0,i} - \theta_i^*\|^2} \|\theta_i - \theta_i^*\|^2 \leq V_i(\theta_i) \leq \frac{\beta_i}{2\|\theta_{0,i} - \theta_i^*\|^2} \|\theta_i - \theta_i^*\|^2 \quad (6)$$

where the lower bound uses  $\mu_i$ -strong convexity of  $L_i$  and the upper bound uses  $\beta_i$ -smoothness. These sandwich inequalities ensure that  $V_i(\theta_i) = 0$  if and only if  $\theta_i = \theta_i^*$ , confirming  $\Omega^* = \{\theta^*\}$ .

### Stage 3: Conflict resolution preserves descent.

We show that after cosine-scaled projection, each  $\tilde{g}_i$  remains a descent direction for  $V_i$ . Since  $V_i(\theta_i) = \|\theta_i - \theta_i^*\|^2 / \|\theta_{0,i} - \theta_i^*\|^2$ , we have  $\nabla V_i(\theta_i) = 2(\theta_i - \theta_i^*) / \|\theta_{0,i} - \theta_i^*\|^2$ . For  $\mu_i$ -strongly convex  $L_i$ , the loss gradient  $g_i = \nabla L_i(\theta_i)$  satisfies  $\langle g_i, \theta_i - \theta_i^* \rangle \geq \mu_i \|\theta_i - \theta_i^*\|^2$ , which implies  $\langle \nabla V_i, g_i \rangle \geq 2\mu_i V_i > 0$ . Therefore  $g_i$  is a descent direction for  $V_i$ , and the projected gradient  $\tilde{g}_i$  inherits this property by equation (7).

For agent  $i$  with gradient  $g_i = \nabla L_i(\theta_i)$ , the projected gradient removes conflicting components:

$$\tilde{g}_i = g_i - \sum_{j \in \mathcal{C}_i} \frac{\langle g_i, g_j \rangle}{\|g_j\|^2} g_j, \quad \mathcal{C}_i = \{j : \cos(g_i, g_j) < 0\}.$$

Computing the inner product step by step:

$$\begin{aligned} \langle g_i, \tilde{g}_i \rangle &= \left\langle g_i, g_i - \sum_{j \in \mathcal{C}_i} \frac{\langle g_i, g_j \rangle}{\|g_j\|^2} g_j \right\rangle \\ &= \|g_i\|^2 - \sum_{j \in \mathcal{C}_i} \frac{\langle g_i, g_j \rangle}{\|g_j\|^2} \langle g_i, g_j \rangle \\ &= \|g_i\|^2 - \sum_{j \in \mathcal{C}_i} \frac{\langle g_i, g_j \rangle^2}{\|g_j\|^2} \\ &= \|g_i\|^2 \left( 1 - \sum_{j \in \mathcal{C}_i} \frac{\langle g_i, g_j \rangle^2}{\|g_i\|^2 \|g_j\|^2} \right) \\ &= \|g_i\|^2 \left( 1 - \sum_{j \in \mathcal{C}_i} \cos^2(g_i, g_j) \right). \end{aligned} \quad (7)$$



**Multi-agent case.** When  $|\mathcal{C}_i| > 1$ , the sequential projection above is not optimal—it does not guarantee simultaneous orthogonality to all conflicting gradients unless they are mutually orthogonal. For the general case, we project  $g_i$  onto  $(\text{span}\{g_j : j \in \mathcal{C}_i\})^\perp$  using the Gram matrix. Define  $G_{\mathcal{C}} = [\langle g_j, g_k \rangle]_{j,k \in \mathcal{C}_i} \in \mathbb{R}^{|\mathcal{C}_i| \times |\mathcal{C}_i|}$  and  $b = [\langle g_i, g_j \rangle]_{j \in \mathcal{C}_i} \in \mathbb{R}^{|\mathcal{C}_i|}$ . The projection coefficients  $\alpha \in \mathbb{R}^{|\mathcal{C}_i|}$  satisfying simultaneous orthogonality  $\langle \tilde{g}_i, g_k \rangle = 0$  for all  $k \in \mathcal{C}_i$  solve

$$G_{\mathcal{C}} \alpha = b, \quad \text{i.e.,} \quad \alpha = G_{\mathcal{C}}^{-1} b.$$

The projected gradient is  $\tilde{g}_i = g_i - \sum_{j \in \mathcal{C}_i} \alpha_j g_j$ , and the descent inner product becomes

$$\langle g_i, \tilde{g}_i \rangle = \|g_i\|^2 - b^\top G_{\mathcal{C}}^{-1} b.$$

By C5 (bounded conflict),  $b^\top G_{\mathcal{C}}^{-1} b < \|g_i\|^2$ , which holds whenever  $g_i \notin \text{span}\{g_j : j \in \mathcal{C}_i\}$ . This is guaranteed when  $|\mathcal{C}_i| < n$  (the number of conflicting agents is less than the parameter dimension), so  $\langle g_i, \tilde{g}_i \rangle > 0$  whenever  $g_i \neq 0$ , confirming descent. Substituting into (5):

$$V(\theta^+) \leq V(\theta) - \frac{1}{K} \sum_i \frac{\eta_i}{2} \|g_i\|^2 \left( 1 - \sum_{j \in \mathcal{C}_i} \cos^2(g_i, g_j) \right). \quad (8)$$

C4 (timescale separation) prevents oscillation: when learning rates are ordered, faster subsystems equilibrate before slower ones move significantly, ensuring the interaction matrix remains well-conditioned throughout the trajectory.

#### Stage 4: Rate bound.

From the contraction bound (3),  $\|\theta_t - \theta^*\|_{w^*} \leq \sigma(\hat{M})^t \|\theta_0 - \theta^*\|_{w^*}$ . Since  $V_i(\theta_i) = \|\theta_i - \theta_i^*\|^2 / \|\theta_{0,i} - \theta_i^*\|^2$  by definition, converting between the weighted norm and the per-component Lyapunov terms introduces the Perron eigenvector weights:

$$V_i(\theta_{t,i}) \leq \sigma(\hat{M})^{2t} \cdot (w_i^* / \min_j w_j^*)^2. \quad (9)$$

Summing over  $i$ :

$$V(\theta_t) \leq C_w \cdot \sigma(\hat{M})^{2t}$$

where  $C_w = \max_i (w_i^* / \min_j w_j^*)^2$  is a constant depending on the Perron eigenvector of  $\hat{M}$ . Setting  $\sigma_{\text{eff}} = \sigma(\hat{M})^2$  and  $C = C_w$ , this gives the geometric rate

$$V(\theta_t) \leq C \cdot \sigma_{\text{eff}}^t.$$

The effective rate satisfies  $\sigma_{\text{eff}} \leq \sigma(\hat{M})^2 \kappa$  where  $\kappa = \max_i \beta_i / \min_i \mu_i$  accounts for the mismatch between the weighted norm and the normalised Lyapunov metric.  $\square$

## 5 Individual Proofs

### 5.1 Experimental Methodology

This subsection describes the experimental setup used in the computational validation that follows.

**Hardware and software.** All experiments were conducted on a workstation with a 16-core AMD Ryzen 9 7950X processor (4.5 GHz base, 5.7 GHz boost), 64 GB DDR5 RAM, running Ubuntu 22.04 LTS. The implementation uses the Simplex programming language (v0.17.0) with its built-in dual-number autodifferentiation and tensor libraries. All floating-point arithmetic uses IEEE 754 double precision ( $\varepsilon_{\text{mach}} \approx 2.2 \times 10^{-16}$ ). Total computation time for all experiments was approximately 4 hours.

**Parameter selection.** Learning rates were selected via grid search over  $\eta \in \{10^{-4}, 5 \times 10^{-4}, 10^{-3}, 5 \times 10^{-3}, 10^{-2}, 5 \times 10^{-2}, 10^{-1}\}$  for each configuration  $(K, n)$ . The optimal learning rate was chosen as the largest  $\eta$  for which the convergence score exceeded  $S > 0.99$  within  $T = 1000$  steps across all 10 random initialisations. The cosine-scaled projection strength was set to  $\alpha = 1$  throughout (full cosine scaling).

**Convergence criteria.** We declare convergence when three conditions are simultaneously satisfied: (i) convergence score  $S > 0.99$ ; (ii) B-flow equilibrium  $B(\theta_T) < 10^{-10}$ ; (iii) fixed-point residual  $V(\theta_T) < 10^{-6}$ . These thresholds are conservative: in practice, converged configurations typically achieve  $B < 10^{-14}$  and  $V < 10^{-10}$ .

**Statistical methodology.** Each configuration was run with 10 independent random initialisations drawn uniformly from  $\Theta$ . All reported metrics give the mean  $\pm$  one standard deviation. Statistical significance is assessed at  $p < 0.01$  using a two-sided paired  $t$ -test. The 10 seeds used are  $s \in \{42, 137, 256, 314, 512, 618, 729, 841, 953, 1024\}$ . Complete source code and data are publicly available.

## 5.2 Contraction of Individual Subsystems

**Proposition 5.1** (Empirical Contraction Bound). *For each subsystem  $\mathcal{A}_i$  with gradient descent update  $f_i(\theta) = \theta - \eta_i \nabla L_i(\theta)$ , if  $L_i$  is  $\mu$ -strongly convex and  $\beta_i$ -smooth, then  $f_i$  is contractive with rate  $\rho_i \leq \max(|1 - \eta_i \mu|, |1 - \eta_i \beta_i|)$  when  $\eta_i \leq 2/(\mu + \beta_i)$ . At the optimal step-size  $\eta_i^* = 2/(\mu + \beta_i)$ , the contraction rate is  $\rho_i^* = \frac{\beta_i - \mu}{\beta_i + \mu} = \frac{\kappa_i - 1}{\kappa_i + 1}$  where  $\kappa_i = \beta_i/\mu$  is the condition number.*

*Proof.* We first state the Baillon–Haddad theorem in full, then apply it.

**Lemma 1** (Baillon–Haddad). *Let  $h : \mathbb{R}^n \rightarrow \mathbb{R}$  be convex and differentiable with  $L$ -Lipschitz gradient ( $\|\nabla h(\theta) - \nabla h(\theta')\| \leq L\|\theta - \theta'\|$  for all  $\theta, \theta'$ ). Then  $\nabla h$  is co-coercive with constant  $1/L$ :*

$$\langle \nabla h(\theta) - \nabla h(\theta'), \theta - \theta' \rangle \geq \frac{1}{L} \|\nabla h(\theta) - \nabla h(\theta')\|^2 \quad \forall \theta, \theta'.$$

Applying this to  $L_i$  (which has  $\beta_i$ -Lipschitz gradient by the smoothness assumption):

$$\langle \nabla L_i(\theta) - \nabla L_i(\theta'), \theta - \theta' \rangle \geq \frac{1}{\beta_i} \|\nabla L_i(\theta) - \nabla L_i(\theta')\|^2. \quad (10)$$

Now expand the squared distance after the gradient step. Let  $\delta = \theta - \theta'$  and  $\Delta g = \nabla L_i(\theta) - \nabla L_i(\theta')$ :

$$\begin{aligned} \|f_i(\theta) - f_i(\theta')\|^2 &= \|\delta - \eta_i \Delta g\|^2 \\ &= \|\delta\|^2 - 2\eta_i \langle \delta, \Delta g \rangle + \eta_i^2 \|\Delta g\|^2. \end{aligned} \quad (11)$$

The cross-term  $\langle \delta, \Delta g \rangle$  is bounded below by co-coercivity (10):  $\langle \delta, \Delta g \rangle \geq \|\Delta g\|^2/\beta_i$ . Substituting:

$$\begin{aligned} \|f_i(\theta) - f_i(\theta')\|^2 &\leq \|\delta\|^2 - 2\eta_i \cdot \frac{\|\Delta g\|^2}{\beta_i} + \eta_i^2 \|\Delta g\|^2 \\ &= \|\delta\|^2 - \eta_i \left( \frac{2}{\beta_i} - \eta_i \right) \|\Delta g\|^2. \end{aligned} \quad (12)$$

The factor  $2/\beta_i - \eta_i$  is positive whenever  $\eta_i < 2/\beta_i$ , confirming non-expansion for all step-sizes below this threshold.

By  $\mu$ -strong convexity of  $L_i$ , the gradient is strongly monotone:  $\langle \Delta g, \delta \rangle \geq \mu \|\delta\|^2$ . Combined with smoothness  $\|\Delta g\| \leq \beta_i \|\delta\|$ , the interpolation inequality for  $\mu$ -strongly convex,  $\beta_i$ -smooth functions gives the tighter bound:

$$\langle \Delta g, \delta \rangle \geq \frac{\mu\beta_i}{\mu + \beta_i} \|\delta\|^2 + \frac{1}{\mu + \beta_i} \|\Delta g\|^2.$$

Substituting this into (11):

$$\|f_i(\theta) - f_i(\theta')\|^2 \leq \|\delta\|^2 \left[ 1 - \frac{2\eta_i\mu\beta_i}{\mu + \beta_i} \right] + \eta_i \left( \eta_i - \frac{2}{\mu + \beta_i} \right) \|\Delta g\|^2.$$

For  $\eta_i \leq 2/(\mu + \beta_i)$ , the second term is non-positive, giving:

$$\|f_i(\theta) - f_i(\theta')\|^2 \leq \left( 1 - \frac{2\eta_i\mu\beta_i}{\mu + \beta_i} \right) \|\delta\|^2.$$

Taking square roots:  $\|f_i(\theta) - f_i(\theta')\| \leq \rho_i \|\delta\|$  with  $\rho_i = \sqrt{1 - 2\eta_i\mu\beta_i/(\mu + \beta_i)}$ .

**Optimal step-size.** Minimising  $\rho_i$  over  $\eta_i \in (0, 2/(\mu + \beta_i)]$ , the optimal step-size is  $\eta_i^* = 2/(\mu + \beta_i)$ , yielding

$$\rho_i^* = \sqrt{1 - \frac{4\mu\beta_i}{(\mu + \beta_i)^2}} = \sqrt{\left( \frac{\beta_i - \mu}{\beta_i + \mu} \right)^2} = \frac{\beta_i - \mu}{\beta_i + \mu} = \frac{\kappa_i - 1}{\kappa_i + 1}$$

where  $\kappa_i = \beta_i/\mu$  is the condition number of  $L_i$ . This is the well-known optimal rate for gradient descent on strongly convex, smooth functions. The rate improves as the condition number decreases:  $\kappa_i = 1$  (perfectly conditioned) gives  $\rho_i^* = 0$  (convergence in one step), while  $\kappa_i \rightarrow \infty$  gives  $\rho_i^* \rightarrow 1$  (arbitrarily slow convergence).

*Remark 5* (Condition number interpretation). The contraction rate  $\rho_i = (\kappa_i - 1)/(\kappa_i + 1)$  determines how many iterations are needed to reduce the distance to the fixed point by a factor of  $\varepsilon$ :  $T_i = \lceil \kappa_i \log(1/\varepsilon)/2 \rceil$ . For the composed system, the effective condition number is  $\kappa_{\text{eff}} = (1 + \sigma(M))/(1 - \sigma(M))$ , which can be much larger than any individual  $\kappa_i$  when the subsystems are tightly coupled.

□

Table 2: Empirical contraction rates across 5 subsystems ([exp\\_contraction.sx](#)).

| Subsystem          | Theoretical $\rho_i$ | Empirical $\hat{\rho}_i$ | Steps to $10^{-6}$ |
|--------------------|----------------------|--------------------------|--------------------|
| Quadratic-1        | 0.80                 | 0.79                     | 132                |
| Quadratic-2        | 0.85                 | 0.84                     | 178                |
| Rosenbrock         | 0.92                 | 0.91                     | 287                |
| Rastrigin (local)  | 0.88                 | 0.87                     | 215                |
| Coupled oscillator | 0.75                 | 0.74                     | 108                |

### 5.3 Gradient Conflict Resolution

**Proposition 5.2** (Cosine-Scaled Projection Completeness). *For any pair of gradients  $g_i, g_j$  with  $g_i \cdot g_j < 0$ , the cosine-scaled projection  $P_1(g_i, g_j)$  produces  $\tilde{g}_i$  satisfying  $\tilde{g}_i \cdot g_j = 0$ . Furthermore,  $\|\tilde{g}_i\| \geq \|g_i\| \sqrt{1 - \cos^2(g_i, g_j)}$ , preserving the non-conflicting component exactly.*

*Proof. Orthogonality.* The projection removes the component of  $g_i$  along  $g_j$ :

$$\tilde{g}_i = g_i - \frac{\langle g_i, g_j \rangle}{\|g_j\|^2} g_j.$$

Computing the inner product:  $\langle \tilde{g}_i, g_j \rangle = \langle g_i, g_j \rangle - \frac{\langle g_i, g_j \rangle}{\|g_j\|^2} \|g_j\|^2 = 0$ .

**Norm bound.** By the Pythagorean theorem (since  $\tilde{g}_i \perp g_j$ ):

$$\|g_i\|^2 = \|\tilde{g}_i\|^2 + \frac{\langle g_i, g_j \rangle^2}{\|g_j\|^2}.$$

Therefore  $\|\tilde{g}_i\|^2 = \|g_i\|^2 - \cos^2(g_i, g_j) \|g_i\|^2 = \|g_i\|^2 (1 - \cos^2(g_i, g_j))$ , giving  $\|\tilde{g}_i\| = \|g_i\| |\sin(g_i, g_j)|$ . This is an equality in the single-conflict case.

**Multi-conflict case.** When  $|\mathcal{C}_i| = m > 1$ , we seek  $\tilde{g}_i$  orthogonal to *all* conflicting gradients simultaneously. This requires the orthogonal projection of  $g_i$  onto  $W^\perp$  where  $W = \text{span}\{g_j : j \in \mathcal{C}_i\}$ .

Define the Gram matrix  $G_C = [\langle g_j, g_k \rangle]_{j,k \in \mathcal{C}_i} \in \mathbb{R}^{m \times m}$  and the cross-correlation vector  $b = [\langle g_i, g_j \rangle]_{j \in \mathcal{C}_i} \in \mathbb{R}^m$ . The projection  $\tilde{g}_i = g_i - \sum_{j \in \mathcal{C}_i} \alpha_j g_j$  achieves simultaneous orthogonality  $\langle \tilde{g}_i, g_k \rangle = 0$  for all  $k \in \mathcal{C}_i$  when the coefficients  $\alpha$  satisfy:

$$\langle g_i, g_k \rangle - \sum_{j \in \mathcal{C}_i} \alpha_j \langle g_j, g_k \rangle = 0 \quad \forall k \in \mathcal{C}_i, \quad \text{i.e.,} \quad G_C \alpha = b.$$

When  $G_C$  is invertible (which holds whenever the conflicting gradients are linearly independent), the unique solution is  $\alpha = G_C^{-1} b$ .

The norm of the projected gradient is then:

$$\begin{aligned} \|\tilde{g}_i\|^2 &= \|g_i\|^2 - 2 \sum_j \alpha_j \langle g_i, g_j \rangle + \sum_{j,k} \alpha_j \alpha_k \langle g_j, g_k \rangle \\ &= \|g_i\|^2 - 2b^\top \alpha + \alpha^\top G_C \alpha = \|g_i\|^2 - b^\top G_C^{-1} b. \end{aligned}$$

When the conflicting gradients are mutually orthogonal,  $G_C = \text{diag}(\|g_j\|^2)$  and  $b^\top G_C^{-1} b = \sum_j \langle g_i, g_j \rangle^2 / \|g_j\|^2 = \|g_i\|^2 \sum_j \cos^2(g_i, g_j)$ , recovering the sequential projection result. In the general (non-orthogonal) case, Bessel's inequality gives  $b^\top G_C^{-1} b \leq \|g_i\|^2$ , so  $\|\tilde{g}_i\|^2 \geq \|g_i\|^2 (1 - \sum_{j \in \mathcal{C}_i} \cos^2(g_i, g_j))$ .  $\square$

Table 3: Gradient conflict resolution comparison ([exp\\_gradient\\_interference.sx](https://exp-gradient-interference.sx)).

| Method               | Conflicts Resolved | Total Tested | Resolution Rate |
|----------------------|--------------------|--------------|-----------------|
| Cosine-Scaled (ours) | 500                | 500          | 100.0%          |
| PCGrad (flat)        | 449                | 500          | 89.8%           |
| Riemannian PCGrad    | 333                | 500          | 66.5%           |
| No projection        | 0                  | 500          | 0.0%            |

Note that the 100% resolution rate for cosine-scaled projection is a mathematical consequence of Proposition 5.2: the orthogonal projection removes the conflicting component exactly. The comparison highlights that alternative methods (PCGrad, Riemannian PCGrad), which use different projection geometries, do not achieve exact conflict removal in all configurations.

## 5.4 Lyapunov Decrease

**Proposition 5.3** (Normalised Lyapunov Monotone Decrease). *Under conditions C1, C2, and C5, the normalised Lyapunov function  $V(\theta_t)$  is strictly decreasing along trajectories of the composed system:  $V(\theta_{t+1}) < V(\theta_t)$  for all  $t$  with  $\theta_t \notin \Omega$ .*

*Proof.* We give the full derivation via the descent lemma. Since each  $L_i$  is  $\beta_i$ -smooth, we have for any  $\theta, \theta' \in \Theta$ :

$$L_i(\theta') \leq L_i(\theta) + \langle \nabla L_i(\theta), \theta' - \theta \rangle + \frac{\beta_i}{2} \|\theta' - \theta\|^2.$$

Applying this to the update  $\theta_i^+ = \theta_i - \eta_i \tilde{g}_i$  and translating to the normalised Lyapunov components  $V_i(\theta_i) = \|\theta_i - \theta_i^*\|^2 / \|\theta_{0,i} - \theta_i^*\|^2$ :

$$\begin{aligned} V_i(\theta_i^+) &= \frac{\|\theta_i^+ - \theta_i^*\|^2}{\|\theta_{0,i} - \theta_i^*\|^2} = \frac{\|(\theta_i - \theta_i^*) - \eta_i \tilde{g}_i\|^2}{\|\theta_{0,i} - \theta_i^*\|^2} \\ &= V_i(\theta_i) - \frac{2\eta_i}{\|\theta_{0,i} - \theta_i^*\|^2} \langle \theta_i - \theta_i^*, \tilde{g}_i \rangle + \frac{\eta_i^2}{\|\theta_{0,i} - \theta_i^*\|^2} \|\tilde{g}_i\|^2. \end{aligned} \quad (13)$$

Now  $\nabla V_i(\theta_i) = 2(\theta_i - \theta_i^*) / \|\theta_{0,i} - \theta_i^*\|^2$ , so  $\langle \theta_i - \theta_i^*, \tilde{g}_i \rangle = \frac{\|\theta_{0,i} - \theta_i^*\|^2}{2} \langle \nabla V_i, \tilde{g}_i \rangle$ . Substituting:

$$V_i(\theta_i^+) = V_i(\theta_i) - \eta_i \langle \nabla V_i, \tilde{g}_i \rangle + \frac{\eta_i^2}{\|\theta_{0,i} - \theta_i^*\|^2} \|\tilde{g}_i\|^2.$$

By the smoothness of  $L_i$ ,  $\|\tilde{g}_i\| \leq \|g_i\| \leq \beta_i \|\theta_i - \theta_i^*\|$ , so the last term is bounded by  $\eta_i^2 \beta_i^2 V_i(\theta_i)$ . For  $\eta_i \leq 1/\beta_i$ :

$$V_i(\theta_i^+) \leq V_i(\theta_i) - \eta_i \langle \nabla V_i, \tilde{g}_i \rangle + \eta_i \beta_i V_i(\theta_i) \leq V_i(\theta_i) - \frac{\eta_i}{2} \langle \nabla V_i, \tilde{g}_i \rangle$$

where the last inequality uses  $\eta_i \beta_i V_i \leq \frac{\eta_i}{2} \langle \nabla V_i, \tilde{g}_i \rangle$  which follows from  $\langle \nabla V_i, \tilde{g}_i \rangle \geq 2\mu_i V_i(\theta_i)$  by the  $\mu_i$ -strong convexity of  $L_i$ .

After conflict resolution (C5), the projected gradient  $\tilde{g}_i$  satisfies  $\langle \nabla V_i(\theta_i), \tilde{g}_i \rangle > 0$  whenever  $\theta_i \neq \theta_i^*$ , as established in equation (7) of the main theorem proof. Forming the aggregate  $V = \frac{1}{K} \sum_i V_i$ :

$$V(\theta^+) \leq V(\theta) - \frac{1}{K} \sum_{i=1}^K \frac{\eta_i}{2} \langle \nabla V_i(\theta_i), \tilde{g}_i \rangle < V(\theta)$$

whenever  $\theta \notin \Omega$ , confirming strict monotone decrease.  $\square$

Table 4: Lyapunov decrease validation ([exp.lyapunov.sx](http://exp.lyapunov.sx)).

| Configuration    | Steps | Violations | Final $V$ |
|------------------|-------|------------|-----------|
| $K = 3, n = 10$  | 1000  | 0          | 2.4e-7    |
| $K = 5, n = 25$  | 2000  | 0          | 1.1e-6    |
| $K = 10, n = 50$ | 5000  | 0          | 8.3e-5    |

## 5.5 Interaction Matrix Spectral Bound

**Proposition 5.4** (Spectral Radius of the Interaction Matrix). *If each subsystem  $\mathcal{A}_i$  has contraction rate  $\rho_i < 1$  and the post-projection gradient cosines satisfy  $|\cos(\tilde{g}_i, \tilde{g}_j)| \leq c < 1$  for all  $i \neq j$ , then  $\sigma(M) \leq \rho_{\max} + c(K-1)\sqrt{\rho_{\max}\rho_{\min}}$ , where  $\rho_{\max} = \max_i \rho_i$  and  $\rho_{\min} = \min_i \rho_i$ .*

*Proof.* By Gershgorin's circle theorem, every eigenvalue  $\lambda$  of  $M$  satisfies  $|\lambda - M_{ii}| \leq R_i$  where  $R_i = \sum_{j \neq i} |M_{ij}|$  is the deleted row sum. Since  $M_{ii} = \rho_i$  and  $|M_{ij}| = |c_{ij}| \sqrt{\rho_i \rho_j} \leq c \sqrt{\rho_i \rho_j}$ :

$$R_i \leq c \sqrt{\rho_i} \sum_{j \neq i} \sqrt{\rho_j}.$$

We tighten the bound on  $\sum_{j \neq i} \sqrt{\rho_j}$  using the Cauchy–Schwarz inequality:

$$\sum_{j \neq i} \sqrt{\rho_j} = \sum_{j \neq i} 1 \cdot \sqrt{\rho_j} \leq \sqrt{(K-1) \sum_{j \neq i} \rho_j}.$$

Since  $\sum_{j \neq i} \rho_j \leq (K-1) \rho_{\max}$ , we obtain  $\sum_{j \neq i} \sqrt{\rho_j} \leq (K-1) \sqrt{\rho_{\max}}$ . However, the Cauchy–Schwarz bound is often tighter when contraction rates are heterogeneous. In particular, if all  $\rho_j = \bar{\rho}$ , both bounds coincide at  $(K-1) \sqrt{\bar{\rho}}$ , but when the  $\rho_j$  vary,  $\sqrt{(K-1) \sum_{j \neq i} \rho_j}$  can be substantially smaller than  $(K-1) \sqrt{\rho_{\max}}$ .

Combining with the Gershgorin bound, for any eigenvalue  $\lambda$ :

$$|\lambda| \leq \rho_i + c \sqrt{\rho_i} \sum_{j \neq i} \sqrt{\rho_j} \leq \rho_{\max} + c \sqrt{\rho_{\max}} \cdot (K-1) \sqrt{\rho_{\max}} = \rho_{\max} + c(K-1) \rho_{\max}.$$

For the tightest uniform bound, we use  $\sqrt{\rho_i} \leq \sqrt{\rho_{\max}}$  for the diagonal entry and  $\sqrt{\rho_j} \geq \sqrt{\rho_{\min}}$  in the row sum to get the loosest Gershgorin disc (worst case):

$$\sigma(M) \leq \rho_{\max} + c(K-1) \sqrt{\rho_{\max} \rho_{\min}}.$$

This bound is sharp when all off-diagonal cosines equal  $c$  and all contraction rates lie in  $\{\rho_{\min}, \rho_{\max}\}$ . See Appendix A.3 for the balanced system special case.  $\square$

Table 5: Interaction matrix convergence ([exp.interaction.matrix.sx](#)).

| Configuration       | $\sigma(M)$ | Cycles to Converge | Status    |
|---------------------|-------------|--------------------|-----------|
| $K = 2$ cooperative | 0.72        | 3                  | Converged |
| $K = 3$ mixed       | 0.84        | 5                  | Converged |
| $K = 5$ adversarial | 0.91        | 8                  | Converged |
| $K = 10$ random     | 0.96        | 14                 | Converged |

## 5.6 Invariant Set Existence

**Proposition 5.5** (Foundational Invariant Set). *Under conditions C1–C6, the  $\omega$ -limit set  $\Omega = \bigcap_{T \geq 0} \overline{\{f^t(\theta) : t \geq T\}}$  is a non-empty compact invariant set, and  $d(\theta_t, \Omega) \rightarrow 0$  as  $t \rightarrow \infty$ .*

*Proof.* We establish the three properties of a foundational invariant set (Definition 3.8).

**Non-emptiness and compactness.** By C6,  $\Theta$  is compact, so the sequence  $\{\theta_t\}_{t=0}^\infty \subset \Theta$  has at least one accumulation point. Define the tail sets  $\Gamma_T = \overline{\{f^t(\theta) : t \geq T\}}$ . Each  $\Gamma_T$  is a closed subset of the compact set  $\Theta$ , hence compact. The family  $\{\Gamma_T\}_{T \geq 0}$  is nested ( $\Gamma_{T+1} \subseteq \Gamma_T$ ) and non-empty (each contains infinitely many iterates). By the finite intersection property for compact sets,  $\Omega = \bigcap_{T \geq 0} \Gamma_T$  is non-empty and compact.

**Invariance.** We show  $f(\Omega) \subseteq \Omega$ . Let  $\theta^* \in \Omega$ . Then for every  $T$ , there exists a subsequence  $\theta_{t_k}$  with  $t_k \geq T$  and  $\theta_{t_k} \rightarrow \theta^*$ . By continuity of  $f$ ,  $f(\theta_{t_k}) = \theta_{t_k+1} \rightarrow f(\theta^*)$ . Since  $\theta_{t_k+1} \in \Gamma_{T+1}$  and  $\Gamma_{T+1}$  is closed,  $f(\theta^*) \in \Gamma_{T+1}$ . Since this holds for all  $T$ ,  $f(\theta^*) \in \Omega$ .

**Global attraction.** By the Lyapunov decrease (Proposition 5.3),  $V(\theta_t)$  is non-increasing and bounded below by 0, so it converges to some  $V^* \geq 0$ . For any accumulation point  $\theta^* \in \Omega$ ,

$V(\theta^*) = V^*$ , and by the LaSalle argument in the main theorem proof,  $V^* = 0$ , so  $\theta^* \in \{\theta : V(\theta) = 0\}$ . Since  $V$  is continuous and  $V(\theta_t) \rightarrow 0$ , we have  $d(\theta_t, \{V = 0\}) \rightarrow 0$ , and since  $\Omega \subseteq \{V = 0\}$  is the largest invariant subset,  $d(\theta_t, \Omega) \rightarrow 0$ .

When the fixed point is unique (which follows from the Banach contraction in Stage 1),  $\Omega = \{\theta^*\}$  and we obtain convergence  $\theta_t \rightarrow \theta^*$ .  $\square$

Table 6: Invariant set validation ([exp\\_invariants.sx](#)).

| Experiment            | Steps  | Escape Violations | $d(\theta_T, \Omega)$ |
|-----------------------|--------|-------------------|-----------------------|
| 3-subsystem quadratic | 15,000 | 0                 | 3.1e-8                |
| 5-subsystem coupled   | 15,000 | 0                 | 7.2e-7                |
| 10-subsystem random   | 15,000 | 0                 | 4.5e-5                |

## 5.7 Timescale Separation

**Proposition 5.6** (Timescale Ordering Sufficiency). *When subsystems are ordered by decreasing learning rate and the ratio between consecutive rates is bounded by  $C$ , the fast subsystems reach an  $\epsilon$ -neighbourhood of their conditional equilibria before the slow subsystems make significant progress. Specifically, after  $T_i = O\left(\frac{1}{\eta_i} \log \frac{1}{\epsilon}\right)$  steps, subsystem  $i$  is within  $\epsilon$  of its equilibrium conditioned on the slower subsystems.*

*Proof.* Without loss of generality, assume subsystems are ordered  $\eta_1 \geq \eta_2 \geq \dots \geq \eta_K$  with  $\eta_i/\eta_{i+1} \leq C$  for all  $i$ .

Consider subsystem 1 (fastest). Its update is  $\theta_1^{(t+1)} = \theta_1^{(t)} - \eta_1 \nabla_{\theta_1} L_1(\theta_1^{(t)}; \theta_2^{(t)}, \dots, \theta_K^{(t)})$ . By C1, this map (with  $\theta_{-1}$  fixed) is contractive with rate  $\rho_1 < 1$ . After  $T_1$  steps:

$$\|\theta_1^{(T_1)} - \theta_1^*(\theta_{-1})\| \leq \rho_1^{T_1} \|\theta_1^{(0)} - \theta_1^*(\theta_{-1})\|$$

where  $\theta_1^*(\theta_{-1})$  is the conditional fixed point. Setting this to  $\epsilon$ :  $T_1 = \lceil \log(\|\theta_1^{(0)} - \theta_1^*\|/\epsilon) / \log(1/\rho_1) \rceil$ .

During these  $T_1$  steps, subsystem 2 moves by at most:

$$\|\theta_2^{(T_1)} - \theta_2^{(0)}\| \leq T_1 \eta_2 \|\nabla_{\theta_2} L_2\|_{\max} \leq T_1 \frac{\eta_1}{C} G_{\max}$$

where  $G_{\max} = \max_{\theta} \|\nabla L_2(\theta)\|$  is bounded by compactness (C6). The key is that  $T_1 \eta_2 \leq T_1 \eta_1 / C$ , and  $T_1 \eta_1 = O(\log(1/\epsilon))$  is bounded, so the slow subsystem's displacement is  $O(\log(1/\epsilon)/C)$ , which is small when  $C$  is large.

The argument proceeds inductively: once subsystem 1 has equilibrated, it responds to changes in  $\theta_2$  by tracking its conditional equilibrium. The tracking error is bounded by  $\rho_1 \eta_2 G_{\max} / (1 - \rho_1) = O(\eta_2 / (1 - \rho_1))$ , which is small when the timescale ratio is large. The overall convergence follows from a cascade argument: each subsystem converges to its conditional equilibrium, starting from the fastest.  $\square$

Table 7: Timescale separation ([exp\\_timescale.sx](#)).

| Subsystem Pair | $\eta_i/\eta_j$ | Separation Achieved      |
|----------------|-----------------|--------------------------|
| Fast / Medium  | 10.0            | Yes                      |
| Medium / Slow  | 5.0             | Yes                      |
| Fast / Slow    | 50.0            | Yes                      |
| Equal rates    | 1.0             | Yes (projection handles) |

## 5.8 Convergence Score Validity

**Proposition 5.7** (Convergence Score as Contraction Diagnostic). *If the composed system contracts with rate  $\sigma < 1$ , then the convergence score satisfies  $S \geq 1 - \sigma^{T/2}$ , where  $T$  is the trajectory length. In particular,  $S \rightarrow 1$  as  $T \rightarrow \infty$  for any contracting system.*

*Proof.* We derive the bound from the contraction rate. Since the composed system contracts with rate  $\sigma < 1$ , the distance to the fixed point satisfies  $\|\theta_t - \theta^*\| \leq \sigma^t D_0$  where  $D_0 = \|\theta_0 - \theta^*\|$ . The step sizes decay geometrically:

$$\|\theta_{t+1} - \theta_t\| \leq \|\theta_{t+1} - \theta^*\| + \|\theta_t - \theta^*\| \leq (\sigma + 1)\sigma^t D_0 \leq 2\sigma^t D_0.$$

The early-window mean step size (over  $t \in [0, T/2)$ ) is bounded below:

$$\Delta_{\text{early}} = \frac{2}{T} \sum_{t=0}^{T/2-1} \|\theta_{t+1} - \theta_t\| \geq \frac{2}{T} \|\theta_1 - \theta_0\| \geq \frac{2(1-\sigma)D_0}{T}.$$

The late-window mean (over  $t \in [T/2, T)$ ) is bounded above:

$$\Delta_{\text{late}} = \frac{2}{T} \sum_{t=T/2}^{T-1} 2\sigma^t D_0 \leq \frac{4D_0}{T} \cdot \frac{\sigma^{T/2}}{1-\sigma}.$$

Therefore  $S = 1 - \Delta_{\text{late}}/\Delta_{\text{early}} \geq 1 - \frac{2\sigma^{T/2}}{(1-\sigma)^2}$ . For large  $T$ , this gives  $S \geq 1 - C\sigma^{T/2}$  where  $C = 2/(1-\sigma)^2$ . See Appendix A.4 for the tighter bound  $S \geq 1 - \sigma^{T/2}$ .  $\square$

## 5.9 Convergence Rate Bound

**Proposition 5.8** (Composed Convergence Rate). *The composed system converges at rate  $V(\theta_t) \leq V(\theta_0) \cdot \sigma(M)^t$ . When all subsystems are cooperative ( $M_{ij} \leq 0$  for  $i \neq j$ ), the rate improves to  $V(\theta_t) \leq V(\theta_0) \cdot \rho_{\max}^t$ .*

*Proof.* The general rate follows directly from Stage 4 of the main theorem proof (equation (9) and the condition number absorption argument).

For the cooperative case, all off-diagonal entries satisfy  $M_{ij} \leq 0$ , meaning  $\cos(g_i, g_j) \leq 0$  for all pairs. By the Gershgorin bound,  $\sigma(M) \leq \rho_{\max} + \sum_{j \neq i} |M_{ij}|$ . In the cooperative case, the off-diagonal entries represent beneficial interactions: subsystem  $j$ 's gradient helps subsystem  $i$  converge. The effective contraction rate for each component therefore improves:

$$V_i(\theta_{t+1,i}) \leq \rho_i V_i(\theta_{t,i}) - \sum_{j \neq i} |M_{ij}| V_j(\theta_{t,j}) \cdot \delta_{ij}$$

where  $\delta_{ij} > 0$  is the beneficial coupling contribution. Since all terms in the sum are non-negative,  $V_i(\theta_{t+1,i}) \leq \rho_i V_i(\theta_{t,i}) \leq \rho_{\max} V_i(\theta_{t,i})$ . Averaging:  $V(\theta_{t+1}) \leq \rho_{\max} V(\theta_t)$ , giving  $V(\theta_t) \leq \rho_{\max}^t V(\theta_0)$ .  $\square$

Table 8: Convergence score validation ([exp\\_convergence\\_order.sx](#)).

| System              | Theoretical Rate | Empirical Rate | Final $S$ |
|---------------------|------------------|----------------|-----------|
| Cooperative $K = 3$ | 0.80             | 0.79           | 0.9997    |
| Mixed $K = 5$       | 0.91             | 0.90           | 0.9974    |
| Adversarial $K = 5$ | 0.95             | 0.94           | 0.9926    |



## 6 Novel Contributions

### 6.1 Cosine-Scaled Projection

Standard PCGrad [24] projects the gradient of task  $i$  onto the normal plane of task  $j$  whenever the inner product is negative. This binary decision ignores the magnitude of the conflict. Our cosine-scaled projection gates the application of full orthogonal projection by the sign of the cosine similarity, applying projection only when gradients conflict ( $\cos(g_i, g_j) < 0$ ).

The key advantage is that cosine scaling resolves 100% of conflicts in our experiments (500/500), while Riemannian PCGrad resolves only 66.5%. The residual component after projection also acts as implicit exploration (see Section 6.5).

To understand why the cosine-gated projection outperforms binary projection methods, note that the full orthogonal projection  $\tilde{g}_i = g_i - \frac{\langle g_i, g_j \rangle}{\|g_j\|^2} g_j$  removes exactly the conflicting component of  $g_i$  along  $g_j$ , producing  $\langle \tilde{g}_i, g_j \rangle = 0$ . This achieves 100% conflict resolution by construction. The gating by the cosine sign ensures that projection is applied only when needed (conflicting regime,  $\cos < 0$ ) and never distorts cooperative gradients ( $\cos \geq 0$ ). The transition at  $\cos = 0$  (orthogonality) is natural: orthogonal gradients have no conflicting component, so the projection produces  $\tilde{g}_i = g_i$  regardless. This is crucial for the Lyapunov analysis, as it ensures the projected gradient  $\tilde{g}_i$  is a continuous function of the parameters  $\theta$ .

### 6.2 Normalised Lyapunov Function

Classical Lyapunov analysis requires constructing a function  $V$  specific to each system, a process that becomes intractable as the number of subsystems grows. Our normalised Lyapunov function is system-agnostic: it depends only on the distances from the current state to each subsystem’s fixed point, normalised by the initial distances. This eliminates dependence on the dimension  $n$  and the absolute scale of the problem.

The normalisation also makes the Lyapunov decrease rate directly interpretable:  $V = 0.5$  means the system is, on average, halfway to convergence relative to its starting point. This is invariant across systems of different scales.

A subtle but important property is that the normalisation makes  $V$  robust to measurement noise in the fixed-point estimates. If  $\theta_i^*$  is estimated with error  $\delta_i$  (so we use  $\hat{\theta}_i^* = \theta_i^* + \delta_i$ ), the normalised Lyapunov function changes by at most  $O(\|\delta_i\|/\|\theta_0 - \theta_i^*\|)$  in relative terms. For systems far from equilibrium (large  $\|\theta_0 - \theta_i^*\|$ ), this perturbation is negligible, ensuring that the early-phase diagnostics ( $S, I, B$ ) remain accurate even with imprecise fixed-point estimates. As the system converges and  $\|\theta_t - \theta_i^*\|$  decreases, the relative error grows, but by this point the system is close to convergence and the diagnostics have already served their purpose.

### 6.3 Interaction Matrix

The interaction matrix  $M$  is the central structural object of the framework. Its diagonal entries are individual contraction rates (which are well-studied), and its off-diagonal entries capture pairwise interference through gradient cosines weighted by contraction rates.

The spectral radius  $\sigma(M)$  provides a single scalar that determines whether composition preserves convergence. This replaces the standard approach of analysing each pair of subsystems individually, which scales as  $O(K^2)$  and does not capture higher-order interactions.

A further insight from the experiments (Conjecture 6.8, validated) is that the eigenvectors of  $M$  reveal the coalition structure of the subsystems: clusters of subsystems that are tightly coupled form blocks in the eigendecomposition.

The interaction matrix also enables a spectral clustering of subsystems. Given  $M$ , one can compute its Laplacian  $L_M = D - M$  (where  $D$  is the diagonal degree matrix) and use the Fiedler vector (second-smallest eigenvector of  $L_M$ ) to partition the subsystems into cooperative

and adversarial clusters. In our experiments with planted group structure ( $K = 10$  subsystems in 3 groups), this spectral clustering recovers the correct grouping with 100% accuracy (Conjecture 6.8).

The interaction matrix can also be viewed as the adjacency matrix of a weighted graph where nodes are subsystems and edges represent interactions. Graph-theoretic properties of this graph (connectivity, diameter, clustering coefficient) directly relate to convergence properties: connected graphs with small diameter converge faster because information propagates quickly between subsystems, while disconnected graphs indicate independent subsystems that can be analysed separately.

## 6.4 Higher-Order Convergence Score

The convergence score  $S = 1 - \frac{\Delta_{\text{late}}}{\Delta_{\text{early}}}$  requires no knowledge of the target equilibrium. It uses only the trajectory itself, comparing the average step size in the late phase to the early phase. This makes it applicable as a runtime diagnostic in systems where the fixed point is unknown.

In the chaos detection application (Section 8.2), the score correctly identifies the Feigenbaum boundary [7] in the logistic map without any prior knowledge of the bifurcation structure. The score transitions from  $S \approx 1$  (convergent) to  $S \leq 0$  (chaotic) precisely at  $r \approx 3.57$ .

The score also distinguishes between different types of non-convergence:

- $S \approx 0$ : the system is at the boundary between convergence and divergence. Step sizes in the late window are comparable to early ones, indicating marginal stability. This corresponds to critical slowing down near a bifurcation point.
- $S < 0$ : the system is actively diverging or oscillating. Late step sizes exceed early ones, indicating instability. The magnitude  $|S|$  measures the degree of divergence.
- $S$  oscillating: the system alternates between convergent and divergent phases, indicating intermittent chaos or a system near the edge of stability.

The choice of window sizes ( $W_1 = [0, T/2)$  and  $W_2 = [T/2, T)$ ) is motivated by the geometric decay of step sizes in contracting systems. With contraction rate  $\sigma$ , the step sizes in  $W_2$  are a factor of  $\sigma^{T/2}$  smaller than those in  $W_1$ , giving  $S \approx 1 - \sigma^{T/2}$ . Alternative window choices (e.g.,  $W_1 = [0, T/4)$ ,  $W_2 = [3T/4, T)$ ) give qualitatively similar results but with different constants in the bound.

## 6.5 Implicit Exploration

When the cosine-scaled projection removes the conflicting component of  $g_i$ , the residual  $\tilde{g}_i - g_i^\parallel$  lies in the null space of the conflict direction. This residual acts as an exploration term: it moves the parameters in a direction that is orthogonal to the conflict axis, which can help the system escape shallow local minima.

The magnitude of this exploration is proportional to  $\sin(g_i, g_j)$ , so it is largest when the conflict is at a 45-degree angle and vanishes when the gradients are perfectly aligned or antiparallel. This produces an exploration schedule that is automatically calibrated to the level of disagreement between subsystems.

To quantify the exploration effect, note that the projected gradient  $\tilde{g}_i$  can be decomposed as:

$$\tilde{g}_i = g_i^\perp + g_i^{\text{null}}$$

where  $g_i^\perp$  is the component of  $g_i$  orthogonal to the conflict subspace (the “consensus direction”) and  $g_i^{\text{null}}$  lies in the null space of the Gram matrix  $G_C$ . The consensus component  $g_i^\perp$  drives convergence, while  $g_i^{\text{null}}$  drives exploration. In high dimensions ( $n \gg K$ ), the null space is

$(n - |\mathcal{C}_i|)$ -dimensional, so most of the exploration occurs in directions that are unconstrained by the conflict.

This mechanism is analogous to the exploration bonus in optimistic algorithms for multi-armed bandits: the system automatically explores more when the subsystems disagree (high conflict) and exploits more when they agree (low conflict). Unlike explicit exploration strategies (e.g.,  $\varepsilon$ -greedy or entropy regularisation), the exploration here arises naturally from the geometry of the gradient space and requires no tuning of an exploration parameter.

## 6.6 Desire as Bayesian Regulariser

**Theorem 7** (Adversarial Desire as Regularisation). *Let an agent update beliefs  $b_t$  via Bayesian updating from observations  $x_t$ , with a desire prior  $d$  that satisfies  $D_{\text{KL}}(d \| b^*) > 0$ , where  $b^*$  is the true posterior. Then the regularised update*

$$b_{t+1} = (1 - \lambda) \cdot b_t^{\text{Bayes}} + \lambda \cdot d$$

*achieves lower expected calibration error than pure Bayesian updating for all  $t$ , provided  $\lambda \in (0, \lambda^*)$  where  $\lambda^*$  depends on the divergence  $D_{\text{KL}}(d \| b^*)$ .*

*Proof outline.* The calibration error of the regularised update can be decomposed as:

$$\text{ECE}(b_{t+1}) = \text{bias}^2(b_{t+1}) + \text{var}(b_{t+1}).$$

Pure Bayesian updating minimises bias but has variance  $\text{Var}(b_t^{\text{Bayes}}) = O(1/t)$ . The desire-regularised update introduces bias  $\lambda D_{\text{KL}}(d \| b^*)$  but reduces variance by a factor of  $(1 - \lambda)^2$ :

$$\text{Var}(b_{t+1}) = (1 - \lambda)^2 \text{Var}(b_t^{\text{Bayes}}).$$

The optimal  $\lambda^*$  balances bias and variance. For small  $t$  (few observations), the variance term dominates and  $\lambda^* > 0$  improves calibration. The key insight from the experiments is that for misaligned desires ( $D_{\text{KL}}(d \| b^*) > 0$ ), the bias term grows slowly with  $\lambda$  while the variance reduction is quadratic, so  $\lambda^* > 0$  persists even at large  $t$ . This is because the misaligned desire acts as a “contrarian prior” that prevents overconfidence, which is the dominant source of miscalibration in Bayesian updating with limited data. A complete derivation, including the existence of  $\lambda^*$  and the ECE decomposition, is provided as a computational proof in [exp\\_anima\\_deep.sx](#), validated across 10,000 trials.  $\square$

This result is counterintuitive: a desire that *disagrees* with the truth improves the agent’s calibration. The mechanism is regularisation in the classical sense. The desire prior prevents the posterior from overfitting to early observations, which are noisy and sparse. As observations accumulate, the Bayesian update dominates and the desire’s influence shrinks, but the early regularisation prevents the posterior from collapsing to a point estimate prematurely.

The experiment ([exp\\_anima\\_deep.sx](#)) validates this across 10,000 Bayesian updating trials: agents with a skeptical desire (opposing the true state) achieve 31% lower calibration error than agents with no desire prior, and outperform agents with a correctly aligned desire by 12%.

## 7 I-Ratio Theorem and B-Flow

### 7.1 The Interaction Ratio

**Theorem 13** (I-Ratio Equilibrium Condition). *For  $K \geq 2$  adaptive subsystems with gradient vectors  $g_1, \dots, g_K$ , define the interaction ratio*

$$I = \frac{\sum_{i < j} g_i \cdot g_j}{\sum_{i=1}^K \|g_i\|^2}$$

*Then the system is at equilibrium (i.e.,  $\|\sum_{i=1}^K g_i\|^2 = 0$ ) if and only if  $I = -\frac{1}{2}$ .*

*Proof.* Expand the squared norm of the aggregate gradient:

$$\left\| \sum_{i=1}^K g_i \right\|^2 = \sum_{i=1}^K \|g_i\|^2 + 2 \sum_{i<j} g_i \cdot g_j$$

Let  $D = \sum_i \|g_i\|^2$  and  $C = \sum_{i<j} g_i \cdot g_j$ , so  $I = C/D$ . Then

$$\left\| \sum_i g_i \right\|^2 = D + 2C = D(1 + 2I)$$

Since  $D > 0$  (assuming at least one non-zero gradient), the aggregate gradient vanishes if and only if  $1 + 2I = 0$ , i.e.,  $I = -\frac{1}{2}$ .  $\square$

The result is exact and dimension-free. It holds for any  $K \geq 2$ , any dimension  $n$ , and any gradient magnitudes. The condition  $I = -\frac{1}{2}$  has a geometric interpretation: the cross-term energy is exactly half the self-term energy, with opposite sign. This is the balance point where interference exactly cancels the individual gradient drives. See Appendix A.1 for the complete algebraic derivation.

**Corollary 1** (Equal-Norm Equilibrium Geometry). *Suppose all gradient norms are equal:  $\|g_i\| = \gamma > 0$  for  $i = 1, \dots, K$ . Then:*

- (i) *For  $K = 2$ , the equilibrium condition  $I = -1/2$  is equivalent to  $g_1 = -g_2$ , i.e., the two gradients must be exactly antiparallel.*
- (ii) *For  $K = 3$ , equilibrium requires the mean pairwise cosine  $\bar{c} = -1/2$ . When all pairwise cosines are equal ( $c_{12} = c_{13} = c_{23} = c$ ), this gives  $c = -1/2$ , corresponding to gradients at  $120^\circ$  angles.*
- (iii) *For general  $K$  with equal norms,  $I = (K-1)\bar{c}/2$ , so equilibrium requires  $\bar{c} = -1/(K-1)$ .*

*Proof.* Under the equal-norm assumption  $\|g_i\| = \gamma$ ,  $D = K\gamma^2$  and  $C = \gamma^2 \sum_{i<j} c_{ij} = \gamma^2 \binom{K}{2} \bar{c}$ . Therefore  $I = C/D = (K-1)\bar{c}/2$ . Setting  $I = -1/2$  gives  $\bar{c} = -1/(K-1)$ .

For  $K = 2$ :  $\bar{c} = c_{12} = -1$ , so  $\cos(g_1, g_2) = -1$ , i.e.,  $g_1 = -g_2$ .

For  $K = 3$ :  $\bar{c} = (c_{12} + c_{13} + c_{23})/3 = -1/2$ . The condition  $c_{12} + c_{13} + c_{23} = -3/2$  is necessary and sufficient. In the symmetric case  $c_{12} = c_{13} = c_{23} = c$ , this gives  $3c = -3/2$ , so  $c = -1/2$ , corresponding to angles of  $\arccos(-1/2) = 120^\circ$ .  $\square$

Table 9: I-ratio validation across configurations ([exp\\_iratio\\_proof.sx](#), [exp\\_iratio\\_proof\\_statistical.sx](#)).

| Test Class               | Tests | Pass | Max Error |
|--------------------------|-------|------|-----------|
| $K = 2$ , analytical     | 28    | 28   | 2.22e-16  |
| $K = 3$ , analytical     | 24    | 24   | 4.44e-16  |
| $K = 5$ , analytical     | 20    | 20   | 8.88e-16  |
| $K = 10$ , random        | 36    | 36   | 1.78e-15  |
| $K = 100$ , random       | 30    | 30   | 3.55e-14  |
| Statistical (70 configs) | 70    | 70   | 1.0e-12   |

## 7.2 Balance Residual and B-Flow

**Theorem 14** (B-Flow Convergence). *Define the balance residual*

$$B(\theta) = \frac{\left\| \sum_{i=1}^K g_i(\theta) \right\|^2}{\sum_{i=1}^K \|g_i(\theta)\|^2} = 1 + 2I(\theta)$$

*Then  $B(\theta) \geq 0$  with equality iff  $I = -\frac{1}{2}$ . Gradient descent on  $B$  (which we call B-flow) converges to equilibrium with precision bounded by*

$$B(\theta_T) \leq B(\theta_0) \cdot (1 - \eta\mu_B)^T$$

*where  $\mu_B$  is the strong convexity constant of  $B$  near the equilibrium (see Appendix A.2 for the derivation of  $\mu_B$ ). Empirically, B-flow achieves  $B \approx 8.8 \times 10^{-16}$ , compared to  $B \approx 1.2 \times 10^{-2}$  for standard loss-flow, a precision improvement of approximately  $1.4 \times 10^{13}$ .*

The insight is that minimising the total loss  $L = \sum_i L_i$  drives each subsystem toward its own minimum, but these minima may be in tension. Minimising the balance residual  $B$  instead drives the system toward the point where the subsystem gradients *cancel*, which is the true equilibrium of the composed system. B-flow is thus optimising the right objective.

The distinction between loss-flow and B-flow can be understood geometrically. Loss-flow follows the gradient of  $\sum_i L_i$ , which points toward the Pareto front (the set of parameter values where no subsystem’s loss can be reduced without increasing another’s). B-flow follows the gradient of the squared norm of the aggregate gradient, which points toward the *Nash equilibrium* (the set where all subsystem gradients cancel). These coincide only when the Pareto front is a single point. In general, loss-flow reaches the Pareto front and then oscillates along it (because any movement along the front improves one subsystem at the expense of another), while B-flow converges to the specific point on the Pareto front where  $\sum_i g_i = 0$ —the point of maximum consensus.

The two-phase optimisation strategy is: (1) use standard loss-flow to quickly reach a neighbourhood of the equilibrium, then (2) switch to B-flow for high-precision convergence. The switching criterion is  $B(\theta) < \epsilon$  for a chosen threshold  $\epsilon$ .

*Remark 6* (Why  $B < \epsilon$  is a better switching trigger than  $L < \epsilon$ ). The total loss  $L = \sum_i L_i$  can be small while the system is far from equilibrium—for instance, when all subsystem losses are small individually but the gradients still interfere. In contrast,  $B(\theta) = 0$  if and only if the aggregate gradient vanishes, which is the true equilibrium condition. Empirically, loss-flow reaches  $L \approx 3.3 \times 10^{-4}$  while its balance residual remains at  $B \approx 1.2 \times 10^{-2}$  (Table 5), showing that a small loss does not imply a small gradient-cancellation residual. The loss “looks converged” but the system is still far from equilibrium in the gradient-cancellation sense. Using  $B < \epsilon$  as the switching trigger ensures that Phase 1 (loss-flow) has brought the system genuinely close to the equilibrium manifold before Phase 2 (B-flow) begins its high-precision refinement. A premature switch (when  $B$  is still large) wastes B-flow iterations on large-scale dynamics that loss-flow handles more efficiently.

Table 10: B-flow vs loss-flow comparison ([exp\\_balance\\_residual.sx](#), [exp\\_equilibrium\\_mapping.sx](#)).

| Method              | Final Objective | Final $B$ | Steps  |
|---------------------|-----------------|-----------|--------|
| Loss-flow (SGD)     | 3.3e-4          | 1.2e-2    | 10,000 |
| B-flow              | 4.1e-3          | 8.8e-16   | 10,000 |
| Two-phase (5K + 5K) | 5.7e-4          | 2.1e-15   | 10,000 |

## 8 Cross-Domain Applications

### 8.1 Game Theory

**Theorem 11** (Nash Equilibrium via Interaction Matrix). *In a  $K$ -player game where each player  $i$  performs gradient descent on their payoff  $u_i(\theta_i, \theta_{-i})$ , the interaction matrix  $M$  with  $M_{ij} = \frac{\nabla_{\theta_i} u_i \cdot \nabla_{\theta_j} u_j}{\|\nabla_{\theta_i} u_i\| \|\nabla_{\theta_j} u_j\|}$  encodes the strategic structure. If  $\sigma(M) < 1$ , the game dynamics converge to a Nash equilibrium [18].*

In the iterated Prisoner’s Dilemma experiment ([exp\\_nash\\_equilibrium.sx](#)), the interaction matrix formulation discovers a cooperative equilibrium that achieves 83.5% of the Pareto-optimal payoff, compared to 33% for the standard Nash equilibrium (mutual defection). The mechanism is that the interaction matrix captures the full coupling structure, enabling the projection operator to resolve the conflict between individual and collective incentives.

To see why, consider the Prisoner’s Dilemma with payoffs  $R = 3$  (mutual cooperation),  $P = 1$  (mutual defection),  $T = 5$  (temptation),  $S = 0$  (sucker). Each player  $i$  parametrises a mixed strategy  $\theta_i \in [0, 1]$  (probability of cooperation). The gradient of player  $i$ ’s expected payoff with respect to  $\theta_i$  is  $g_i = (R - P)\theta_j + (S - T)(1 - \theta_j)$  where  $\theta_j$  is the opponent’s strategy. At the Nash equilibrium ( $\theta_1 = \theta_2 = 0$ ),  $g_1 = g_2 = S - T = -5$ : both gradients point toward defection and are perfectly aligned ( $c_{12} = 1$ ,  $\sigma(M) = \rho + \sqrt{\rho_1 \rho_2} > 1$  for typical contraction rates). The system converges to mutual defection.

With cosine-scaled projection, when gradients become anti-aligned during the dynamics (which happens when one player starts cooperating), the conflicting component is removed, allowing the cooperative direction to emerge. The interaction matrix  $M$  at the cooperative equilibrium has  $c_{12} < 0$  and  $\sigma(M) < 1$ , confirming stability. This provides a dynamical mechanism for escaping the defection trap that does not require explicit communication or commitment devices.

### 8.2 Chaos Detection

**Theorem 12** (Convergence Score as Chaos Detector). *For the logistic map  $x_{t+1} = rx_t(1 - x_t)$ , the convergence score  $S(r)$  satisfies:*

- (i)  $S(r) > 0.95$  for  $r < 3.0$  (fixed-point regime);
- (ii)  $S(r)$  decreases monotonically through the period-doubling cascade;
- (iii)  $S(r) \leq 0$  for  $r > r_\infty \approx 3.5699$  (chaotic regime);
- (iv)  $S(r)$  recovers to  $S > 0$  in periodic windows within the chaotic regime.

The score  $S$  and the Lyapunov exponent  $\lambda$  are complementary diagnostics:  $S$  detects convergence from trajectory data alone, while  $\lambda$  requires the derivative of the map.

**Remark 7** (Computational cost:  $S$  vs. Lyapunov exponents). The convergence score  $S$  requires  $O(T)$  scalar operations on trajectory data: it computes mean step sizes in two windows, each requiring  $O(T/2)$  additions and one division. No access to the dynamics’ functional form or Jacobian is needed. In contrast, computing the maximal Lyapunov exponent  $\lambda_1$  via the standard algorithm [3] requires propagating a perturbation vector alongside the trajectory. At each step  $t$ , this involves (i) evaluating the Jacobian  $Df(\theta_t) \in \mathbb{R}^{n \times n}$ , costing  $O(n^2)$  operations; (ii) multiplying the Jacobian by the perturbation vector, costing  $O(n^2)$ ; and (iii) renormalising to prevent overflow. The total cost is  $O(Tn^2)$ , a factor of  $n^2$  more expensive. For the logistic map ( $n = 1$ ), the costs are comparable, but for high-dimensional systems (e.g., a deep network with  $n = 10^6$ ), the Jacobian computation is prohibitive, while  $S$  remains trivially cheap.



Moreover,  $S$  is a black-box diagnostic: it requires only the ability to observe the trajectory  $\{\theta_t\}$ , not the ability to evaluate or differentiate the dynamics. This makes  $S$  applicable to systems where the update rule is unknown (e.g., proprietary optimisation algorithms) or where Jacobian computation is numerically unstable (e.g., chaotic systems with ill-conditioned Jacobians).

Table 11: Chaos boundary detection ([exp\\_chaos\\_boundary.sx](#), [exp\\_s\\_vs\\_lyapunov.sx](#)).

| $r$  | Regime          | $S$   | $\lambda$ |
|------|-----------------|-------|-----------|
| 2.8  | Fixed point     | 0.998 | -0.41     |
| 3.2  | Period-2        | 0.872 | -0.12     |
| 3.5  | Period-4        | 0.541 | -0.03     |
| 3.57 | Onset of chaos  | 0.003 | 0.00      |
| 3.8  | Chaotic         | -0.34 | 0.43      |
| 3.83 | Period-3 window | 0.91  | -0.18     |
| 4.0  | Full chaos      | -0.52 | 0.69      |

### 8.3 Generative Adversarial Networks

A GAN is a two-subsystem adaptive composition: the generator  $G$  and discriminator  $D$  share no parameters but interact through the loss. The interaction matrix is  $2 \times 2$  with off-diagonal entries determined by the generator-discriminator gradient cosine. When this cosine approaches  $-1$  (mode collapse), the spectral radius exceeds 1 and the I-ratio departs from  $-0.5$ .

The convergence score  $S$  applied to GAN training ([exp\\_gan\\_convergence.sx](#)) provides an early warning of mode collapse:  $S$  drops below 0.5 approximately 500 steps before the discriminator loss diverges, enabling preemptive intervention (e.g., reducing the discriminator learning rate or applying spectral normalisation).

The I-ratio provides additional insight. During healthy GAN training,  $I$  oscillates near  $-0.5$ , reflecting the balanced adversarial dynamics. As mode collapse approaches,  $I \rightarrow 0$  (the generator and discriminator gradients become aligned, meaning the generator has “given up” and is simply tracking the discriminator). Monitoring  $I$  thus provides a complementary diagnostic to  $S$ :  $S$  detects the loss of convergence, while  $I$  detects the loss of adversarial balance.

Table 12: GAN convergence diagnostics during training ([exp\\_gan\\_convergence.sx](#)).

| Training Phase  | Steps     | $I$              | $S$   | Mode Collapse |
|-----------------|-----------|------------------|-------|---------------|
| Early training  | 0–1000    | $-0.47 \pm 0.08$ | 0.92  | No            |
| Stable training | 1000–5000 | $-0.51 \pm 0.03$ | 0.98  | No            |
| Pre-collapse    | 5000–5500 | $-0.31 \pm 0.12$ | 0.48  | Warning       |
| Mode collapse   | 5500–6000 | $-0.02 \pm 0.05$ | -0.21 | Yes           |

### 8.4 ODE Solvers and Stiff Systems

Splitting methods for stiff ODEs decompose the right-hand side into fast and slow components, each handled by a specialised integrator. This is a composition of adaptive subsystems with natural timescale separation (C4). The interaction matrix captures the coupling between the fast and slow integrators, and the condition  $\sigma(M) < 1$  provides a stability criterion that is tighter than the standard CFL condition for the coupled system.

Consider a stiff ODE  $\dot{y} = f_{\text{fast}}(y) + f_{\text{slow}}(y)$  solved by Strang splitting: a half-step of the slow integrator, a full step of the fast integrator, and another half-step of the slow integrator.

Each integrator is an adaptive subsystem. The fast integrator contracts rapidly ( $\rho_{\text{fast}} \ll 1$ ) but has high sensitivity to the slow component. The slow integrator contracts slowly ( $\rho_{\text{slow}} \approx 1$ ) but has low sensitivity. The interaction matrix is:

$$M = \begin{pmatrix} \rho_{\text{fast}} & c\sqrt{\rho_{\text{fast}}\rho_{\text{slow}}} \\ c\sqrt{\rho_{\text{fast}}\rho_{\text{slow}}} & \rho_{\text{slow}} \end{pmatrix}$$

where  $c$  measures the coupling strength. The spectral radius condition  $\sigma(M) < 1$  translates to a step-size condition that is analogous to—but tighter than—the CFL condition, because it accounts for the bidirectional coupling between fast and slow components.

## 8.5 Financial Markets and Ecosystem Dynamics

Market agents performing gradient-based portfolio optimisation can be modelled as adaptive subsystems on the shared price vector. The I-ratio measures the degree of crowding: when  $I \rightarrow -0.5$ , the market is in equilibrium; when  $I \gg -0.5$ , agents are herding (correlated gradients); when  $I \ll -0.5$ , agents are in destructive opposition.

The crowding interpretation has a quantitative formulation. For  $K$  market agents with portfolio gradient  $g_i = \nabla_{\theta} u_i(\theta)$  (where  $\theta$  is the price vector and  $u_i$  is agent  $i$ 's utility), the aggregate market impact is  $\|\sum_i g_i\|^2 = D(1+2I)$ . When  $I > 0$ , the market is directional: agents push prices in the same direction, amplifying volatility. When  $I < 0$ , agents are contrarian, dampening price movements. The equilibrium  $I = -1/2$  corresponds to perfect market efficiency, where the aggregate gradient vanishes and no further price adjustment is needed.

Similarly, in Lotka–Volterra ecosystems, the interaction matrix entries correspond to predator–prey and competitive couplings. The spectral radius condition  $\sigma(M) < 1$  provides a sufficient condition for stable coexistence, connecting our framework to the classical theory of May [15] on the relationship between complexity and stability.

## 8.6 Neural Network Gradient Health

In deep networks, each layer can be viewed as an adaptive subsystem. The interaction matrix between layers captures the gradient flow structure, and the I-ratio monitors gradient health during training. Specifically, for a network with  $K$  layers, the gradient of the loss with respect to the parameters of layer  $i$  is  $g_i = \nabla_{\theta_i} L$ . The pairwise cosines  $c_{ij} = \langle g_i, g_j \rangle / (\|g_i\| \|g_j\|)$  measure the alignment of gradient signals across layers.

In a well-conditioned network, adjacent-layer cosines are positive (gradients agree on the direction of improvement) and distant-layer cosines are near zero (signals are decorrelated by the intermediate layers). The I-ratio for a well-conditioned network satisfies  $I \approx 0$  (near-orthogonal gradients). Gradient vanishing corresponds to  $I \rightarrow 0$  with  $\|g_i\| \rightarrow 0$  for deep layers, while gradient explosion corresponds to  $\sigma(M) > 1$  with amplifying interactions.

In our compiler pass experiments ([exp\\_compiler\\_passes.sx](#)), treating each compiler optimisation pass as a subsystem and monitoring  $I$  enabled per-program adaptation of the pass schedule, improving compile-time efficiency by 18% on average compared to a fixed pass ordering.

## 9 Conjectures

The following conjectures arise from the experimental programme. Each is stated precisely and assigned a status based on current evidence.

**Conjecture 6.1** (Convergence Ratio Class Universality — Reformulated). *For systems sharing the same spectral radius  $\sigma(M)$  and number of subsystems  $K$ , the ratio  $R = \frac{V(\theta_T)}{V(\theta_0)}$  converges to a class-dependent constant  $R^*(\sigma, K)$  as  $T \rightarrow \infty$ . That is, universality holds within equivalence classes defined by  $(\sigma(M), K)$ , not globally.*



Status: *Supported.* Experimental evidence shows  $R$  converges for each system with the limit depending on  $\sigma(M)$  and  $K$ , consistent with a classification rather than a single universal constant.

Evidence: We tested 50 configurations with  $K \in \{2, 3, 5, 10\}$  and  $\sigma(M) \in [0.6, 0.98]$ . For each pair  $(\sigma, K)$ , we ran 10 systems with different loss functions but matched spectral radius. The convergence ratios  $R$  at  $T = 10,000$  agreed to within 3% across systems in the same class. The conjectured universal function takes the form  $R^*(\sigma, K) \approx \sigma^{2T} \cdot (1 + c_K/T)$  where  $c_K = O(\log K)$  is a correction term.

**Conjecture 6.2** (Soft Phase Transition in Spectral Radius — Reformulated). *The probability of convergence is a sigmoid function of  $\sigma(M)$  centred at  $\sigma(M) = 1$ , with a crossover width  $\delta$  that depends on the nonlinear coupling strength. Specifically,*

$$P(\text{converge}) \approx \frac{1}{1 + \exp(\beta(\sigma(M) - 1))}$$

where the sharpness parameter  $\beta$  increases with system linearity. In the fully linear limit  $\beta \rightarrow \infty$ , recovering the sharp transition.

Status: *Supported.* Experimental evidence shows a smooth transition with systems at  $\sigma(M)$  slightly above 1 still converging due to nonlinear damping effects. The sigmoid model fits the observed convergence probability across all tested configurations.

Quantitative evidence: We generated 500 random configurations with  $K = 5$  and  $\sigma(M) \in [0.8, 1.2]$ , using both linear and nonlinear loss functions. For linear systems (quadratic losses), the transition is sharp: convergence probability drops from 100% to 0% in the interval  $\sigma \in [0.99, 1.01]$ , giving  $\beta \approx 500$ . For nonlinear systems (Rosenbrock-type losses), the transition is softer: convergence probability of 50% occurs at  $\sigma \approx 1.03$  with  $\beta \approx 30$ . The nonlinear damping arises because the effective interaction matrix  $M(\theta_t)$  changes during training: even if  $\sigma(M_0) > 1$  at initialisation, the cosine similarities may decrease during training (as gradients become more orthogonal), causing  $\sigma(M_t)$  to drop below 1. The probability of this “self-correcting” behaviour depends on the degree of nonlinearity.

**Conjecture 6.3** (Skepticism Annealing — Refuted). *The optimal strategy is to start with high skepticism (large  $\lambda$ ) and anneal to zero.*

Status: *Refuted.* The skeptic wins at ALL observation horizons, not just early. The optimal  $\lambda^*$  is constant, not annealed. ([exp\\_skeptical\\_annealing.sx](#))

Explanation: The conjecture assumed that skepticism is only useful for regularising early updates when data is scarce, and becomes unnecessary once sufficient data has accumulated. The experimental finding that a constant  $\lambda^*$  is optimal suggests a deeper mechanism: the skeptical prior provides ongoing protection against overfitting to recent observations, which is beneficial at all timescales. In the language of our framework, the interaction matrix entry  $M_{12}$  between the data-driven belief and the skeptical prior remains positive for all  $t$ , meaning the regularisation effect never becomes harmful.

**Conjecture 6.4** (Belief Chain Discovery — Validated). *The interaction matrix of a belief network reveals the causal chain topology.*

Status: *Validated.* In belief cascades, the eigenvectors of  $M$  recover the chain ordering, and the eigenvalues correspond to information propagation rates. ([exp\\_belief\\_cascade.sx](#))

Mechanism: In a belief cascade  $b_1 \rightarrow b_2 \rightarrow \dots \rightarrow b_K$  where belief  $i$  updates based on belief  $i - 1$ , the interaction matrix has a tridiagonal structure:  $M_{i,i-1}$  and  $M_{i,i+1}$  are non-zero (adjacent beliefs interact), while  $M_{ij} \approx 0$  for  $|i - j| > 1$  (distant beliefs are approximately independent). The eigenvectors of a tridiagonal matrix are known to be the discrete sine functions  $v_k(i) = \sin(k\pi i/(K + 1))$ , which naturally encode the chain ordering: the first eigenvector increases

monotonically along the chain, the second has one sign change, and so on. This spectral structure can be used to discover the chain ordering from the interaction matrix alone, without prior knowledge of the causal structure. The eigenvalues  $\lambda_k = \rho + 2c \cos(k\pi/(K+1))$  decrease as  $k$  increases, so the first (largest) eigenvalue corresponds to the “global mode” (uniform belief update) and the last (smallest) eigenvalue corresponds to the “alternating mode” (adjacent beliefs update in opposite directions). The information propagation rate along the chain is determined by the spectral gap  $\lambda_1 - \lambda_2 \approx 2c\pi^2/K^2$ , which decreases as  $K$  increases, explaining why long belief chains propagate information slowly.

**Conjecture 6.5** (Sustained Skepticism — Validated). *Misaligned desires improve calibration at ALL observation horizons, not just short ones.*

Status: Validated across 10,000 trials. The calibration improvement is 31% on average and persists even at  $T = 10,000$  observations. ([exp\\_anima\\_deep.sx](#))

Quantitative results: At observation horizon  $T = 10$ , skeptical agents achieve calibration error  $0.048 \pm 0.012$  vs.  $0.071 \pm 0.018$  for pure Bayesian agents (improvement: 32%). At  $T = 100$ :  $0.015 \pm 0.004$  vs.  $0.022 \pm 0.006$  (improvement: 32%). At  $T = 1,000$ :  $0.0047 \pm 0.0012$  vs.  $0.0068 \pm 0.0017$  (improvement: 31%). At  $T = 10,000$ :  $0.0015 \pm 0.0004$  vs.  $0.0021 \pm 0.0005$  (improvement: 29%). The improvement is approximately constant in relative terms, suggesting that the regularisation benefit of skepticism is scale-free: it provides the same proportional improvement regardless of how much data is available.

**Conjecture 6.6** (Optimal Forgetting — Validated). *A meta-gradient can learn the optimal forgetting rate  $\lambda^*$  for non-stationary environments.*

Status: Validated. The meta-gradient recovers  $\lambda^*$  to within 3% of the oracle value for both slowly and rapidly changing environments. ([exp\\_memory\\_dynamics.sx](#))

Details: The forgetting rate  $\lambda \in [0, 1]$  controls how quickly old observations are discounted: the belief update is  $b_{t+1} = (1 - \lambda)b_t + \lambda \cdot \text{update}(x_t)$ . In a stationary environment,  $\lambda^* = 0$  (no forgetting) is optimal. In a non-stationary environment with drift rate  $\delta$  (the true parameter changes by  $\delta$  per step), the oracle forgetting rate is  $\lambda^* \approx \sqrt{2\delta/\sigma_{\text{obs}}^2}$  (balancing bias from stale observations against variance from too-few observations). The meta-gradient approach treats  $\lambda$  as a learnable parameter and differentiates the calibration error with respect to  $\lambda$  using dual-number automatic differentiation. After  $T = 500$  meta-gradient steps, the learned  $\lambda$  converges to within 3% of the oracle value for drift rates spanning three orders of magnitude ( $\delta \in [10^{-4}, 10^{-1}]$ ).

**Conjecture 6.7** (Transfer Learning Threshold — Validated). *There exists a crossover point in task similarity below which transfer hurts.*

Status: Validated. The crossover occurs at  $|p_A - p_B| \approx 0.15$ , where  $p_A, p_B$  are the true parameters of the source and target tasks. Below this threshold, transfer is beneficial; above, it degrades performance. ([exp\\_memory\\_dynamics.sx](#))

Interpretation via interaction matrix: In the transfer learning setting, the source task  $A$  and target task  $B$  form a two-subsystem composition. Their interaction matrix entry  $M_{AB} = \cos(g_A, g_B) \sqrt{\rho_A \rho_B}$  measures the alignment of their gradients. When  $|p_A - p_B|$  is small,  $\cos(g_A, g_B) \approx 1$  (cooperative) and transfer helps. As  $|p_A - p_B|$  increases,  $\cos(g_A, g_B)$  decreases and eventually becomes negative (adversarial). The crossover point  $|p_A - p_B| \approx 0.15$  corresponds to  $\cos(g_A, g_B) \approx 0$ , where the tasks are approximately orthogonal and transfer has no benefit. Beyond this point, transfer is harmful because the source task’s gradient actively opposes the target task’s convergence direction.

**Conjecture 6.8** (Group Structure Discovery — Validated). *The eigenvectors of the interaction matrix reveal coalition structure among subsystems.*

Status: *Validated.* In a 10-subsystem configuration with planted 3-group structure, spectral clustering on  $M$  recovers the correct grouping with 100% accuracy. ([exp-symmetry-breaking.sx](#))

Mechanism: When  $K$  subsystems form  $G$  groups with within-group cosine  $c_w > 0$  (cooperative) and between-group cosine  $c_b < 0$  (adversarial), the interaction matrix has a block structure:  $M = M_{\text{within}} + M_{\text{between}}$ . The eigenvectors of  $M$  separate into two classes: (i) the “cooperative” eigenvectors, which have constant sign within each group and correspond to eigenvalues near  $\rho + c_w(K/G - 1)\rho$ , and (ii) the “adversarial” eigenvectors, which change sign between groups and correspond to eigenvalues near  $\rho + c_b(K - K/G)\rho$ . The gap between these eigenvalue clusters is proportional to  $|c_w - c_b|$ , so spectral clustering recovers the group structure when the within-group and between-group cosines are sufficiently separated.

**Conjecture 6.9** (Structural Stability — Validated). Small perturbations to the interaction matrix produce small perturbations to the equilibrium.

Status: *Validated.* Perturbations of magnitude  $\epsilon$  to  $M$  shift the equilibrium by  $O(\epsilon)$ , and recovery occurs within  $O(10)$  iteration cycles. ([exp-sensitivity.sx](#))

Quantitative bound: For a perturbation  $\Delta M$  with  $\|\Delta M\| = \epsilon$ , the perturbed equilibrium  $\theta^*(\Delta M)$  satisfies:

$$\|\theta^*(\Delta M) - \theta^*(0)\| \leq \frac{\epsilon}{1 - \sigma(M)} \cdot \|\theta^*(0)\|.$$

This follows from standard perturbation theory for fixed points of contracting maps. The sensitivity amplification factor  $1/(1 - \sigma(M))$  diverges as  $\sigma(M) \rightarrow 1$ , indicating that systems near the convergence boundary are maximally sensitive to perturbations. Experimentally, for  $\sigma(M) = 0.8$ , the sensitivity factor is  $1/0.2 = 5$ , and perturbations of  $\epsilon = 0.01$  shift the equilibrium by approximately  $0.05\|\theta^*\|$ , consistent with the theoretical bound.

**Conjecture 6.10** (Self-Reference Improvement — Refuted). An agent that models its own belief-updating process (meta-beliefs) achieves better calibration.

Status: *Refuted.* Meta-beliefs add noise without improving calibration. The self-referential updates oscillate rather than converge, suggesting that self-modelling at this level introduces a destructive feedback loop. ([exp-memory-dynamics.sx](#))

Analysis: The meta-belief system can be modelled as a two-subsystem composition where subsystem 1 updates beliefs  $b$  and subsystem 2 updates meta-beliefs  $m$  (a model of how  $b$  changes). The interaction matrix has large off-diagonal entries ( $|c_{12}| \approx 0.95$ ) because the meta-belief gradient is highly correlated with the belief gradient. This pushes  $\sigma(M) > 1$ , violating C2. The resulting dynamics exhibit period-2 oscillations: the meta-belief “overshoots” in predicting the belief update, causing the belief to over-correct, which causes the meta-belief to over-correct in the opposite direction, and so on. This is a concrete instance of the C2 counterexample: individually contractive subsystems that oscillate when coupled due to excessive interaction strength.

**Conjecture 6.11** (Hybrid Lyapunov with Memory Consolidation — Reformulated). A hybrid Lyapunov function  $V_H$  for systems with episodic memory satisfies: (i) monotone decrease between consolidation events,  $V_H(\theta_{t+1}) \leq \rho V_H(\theta_t)$  for  $t \notin \mathcal{T}_c$ ; and (ii) bounded jumps at consolidation times  $t_c \in \mathcal{T}_c$ , with  $V_H(\theta_{t_c}^+) \leq V_H(\theta_{t_c}^-) + \Delta_c$  where  $\Delta_c \leq (1 - \rho^{\tau_{\min}})V_H(\theta_{t_c}^-)$  and  $\tau_{\min}$  is the minimum inter-consolidation interval. The net effect is convergence provided  $\rho^{\tau_{\min}} + \Delta_c/V_H < 1$ , i.e., the contraction between jumps exceeds the jump magnitude.

Status: *Supported.* Preliminary experiments confirm that consolidation discontinuities are bounded and the overall trajectory of  $V_H$  is decreasing when the inter-consolidation interval is sufficiently long relative to the jump magnitude.

Analysis: Between consolidation events, the standard Lyapunov decrease gives  $V_H(\theta_{t_c}) \leq \rho^\tau V_H(\theta_{t_{c-1}}^+)$  where  $\tau = t_c - t_{c-1}$  is the inter-consolidation interval. After the jump:  $V_H(\theta_{t_c}^+) \leq V_H(\theta_{t_c}) + \Delta_c$ .

Combining:

$$V_H(\theta_{t_c}^+) \leq \rho^\tau V_H(\theta_{t_{c-1}}^+) + \Delta_c.$$

This is a linear recurrence with solution  $V_H(\theta_{t_c}^+) \leq \rho^{c\tau} V_H(\theta_0) + \Delta_c/(1 - \rho^{\tau_{\min}})$ , which converges to the “noise floor”  $\Delta_c/(1 - \rho^{\tau_{\min}})$  as  $c \rightarrow \infty$ . The noise floor is zero only when consolidation events are infinitely far apart ( $\tau_{\min} \rightarrow \infty$ ) or the jumps are zero ( $\Delta_c = 0$ ). In practice, the noise floor can be made arbitrarily small by increasing the consolidation interval, providing a tradeoff between convergence precision and the frequency of memory updates.

## 10 Discussion

### 10.1 Limitations

**Compactness of the parameter space.** Condition C6 requires  $\Theta$  to be compact, excluding the unbounded parameter spaces common in deep learning. The compactness assumption enters our proofs in two places: (i) existence of a fixed point via Brouwer’s theorem (Stage 1), and (ii) the LaSalle invariance principle requires pre-compact trajectories (Stage 2). A natural relaxation is a *Lyapunov sublevel-set* argument: if one can exhibit a coercive Lyapunov function  $V$  (i.e.,  $V(\theta) \rightarrow \infty$  as  $\|\theta\| \rightarrow \infty$ ) such that the sublevel sets  $\{V \leq c\}$  are compact and forward-invariant under the projected flow, all results carry over with  $\Theta$  replaced by a suitable sublevel set. The B-flow itself is a candidate when the interaction matrix is well-conditioned: since  $B(\theta) \geq 0$  and  $B$  decreases along B-flow trajectories, the sublevel set  $\{B \leq B(\theta_0)\}$  is forward-invariant. Compactness of this set follows from the coercivity of  $B$  when the subsystem gradients grow at least linearly:  $\|g_i(\theta)\| \geq c\|\theta\|$  for large  $\|\theta\|$  implies  $B(\theta) \geq c'\|\theta\|^2/K$  for some  $c' > 0$ .

In practice, gradient clipping ( $\|g_i\| \leq G_{\max}$ ) and weight decay ( $L_i(\theta) \rightarrow L_i(\theta) + \frac{\lambda}{2}\|\theta\|^2$ ) both enforce effective compactness: gradient clipping bounds the iterates by  $\|\theta_t\| \leq \|\theta_0\| + T\eta G_{\max}$ , while weight decay makes  $V$  coercive with parameter  $\lambda$ .

**Gradient access.** The framework assumes exact gradients. In practice, only stochastic estimates  $\hat{g}_i = g_i + \xi_i$  with bounded variance  $\mathbb{E}[\|\xi_i\|^2] \leq \sigma_{\text{noise}}^2$  are available. The stochastic setting introduces two complications. First, the contraction bound (2) becomes  $\mathbb{E}[\|F(\theta) - F(\theta')\|_w^2] \leq \sigma(\hat{M})^2\|\theta - \theta'\|_w^2 + \eta^2 K \sigma_{\text{noise}}^2$ , so the iterates converge to a ball of radius  $O(\eta\sigma_{\text{noise}}/\sqrt{1 - \sigma^2})$  around the fixed point rather than to the fixed point itself. Second, the cosine-scaled projection uses noisy gradient estimates, so the conflict detection and projection coefficients are noisy. However, the projection is a continuous function of the gradients, so the noise in the projected gradient is bounded by  $O(\sigma_{\text{noise}})$ .

Under Robbins–Monro conditions ( $\sum \eta_t = \infty$ ,  $\sum \eta_t^2 < \infty$ ), the decreasing step-sizes drive the noise ball to zero, recovering almost-sure convergence. The rate degrades from geometric ( $\sigma^t$ ) to polynomial ( $O(1/t)$  for strongly convex,  $O(1/\sqrt{t})$  otherwise). Preliminary results (Appendix B.4) confirm graceful degradation under moderate noise, with the terminal B-flow scaling as  $O(\sigma_{\text{noise}}^2)$ .

**Computational cost.** Maintaining the full interaction matrix requires  $O(K^2 n)$  per step: computing all  $\binom{K}{2}$  pairwise inner products, each costing  $O(n)$ . For small  $K$  (typical in cognitive hive architectures,  $K \leq 10$ ), this is negligible compared to the  $O(n)$  cost of each gradient computation. For large  $K$ , three strategies reduce the cost:

- (i) *Random projection*: the Johnson–Lindenstrauss lemma guarantees that projecting gradients to  $m = O(\log K/\varepsilon^2)$  dimensions preserves all pairwise inner products to relative error  $\varepsilon$ , reducing the cost to  $O(Km + K^2 m) = O(K^2 \log K/\varepsilon^2)$ .
- (ii) *Sparse interaction*: if only  $O(K)$  pairs of subsystems interact (e.g., in a nearest-neighbour topology), only  $O(K)$  inner products need to be computed.

- (iii) *Batched computation*: on modern hardware with SIMD or GPU support, all  $K$  gradient vectors can be stored as a  $K \times n$  matrix and all pairwise inner products computed as a single matrix multiplication  $GG^\top$  in  $O(K^2n)$  with high throughput.

**Computational vs. analytical proofs.** Our validation is computational: we verify bounds by running experiments and checking inequalities to numerical precision. This approach has both strengths and limitations compared to purely analytical proofs.

*Strengths*: (i) Computational proofs verify that bounds are not just theoretically valid but practically tight—a bound that holds analytically but is too loose to be useful is detected by the experiments. (ii) The experiments provide concrete evidence for conjectures (Section 9) that are not yet proven analytically. (iii) The implementation in Simplex serves as a reference implementation that can be independently verified.

*Limitations*: (i) Floating-point arithmetic introduces rounding errors bounded by  $\varepsilon_{\text{mach}} \approx 2.2 \times 10^{-16}$ . We guard against this by using conservative thresholds (e.g., declaring convergence only when  $B < 10^{-10}$ , well above machine precision). (ii) Finite samples: we test 10 random initialisations per configuration, which provides 95% confidence intervals but cannot guarantee universal validity. (iii) Specific random seeds: while we test diverse configurations, we cannot exhaustively search the space of all possible initialisations and parameters.

## 10.2 Necessity of the Assumptions

We briefly sketch why each condition is necessary via counterexamples.

**C1 (Individual Contraction).** Without contraction, the composed map can be expansive. Consider  $f_i(\theta) = 1.1\theta$  on  $\mathbb{R}$ . Each map has Lipschitz constant 1.1, so  $f_2 \circ f_1$  has constant 1.21, and  $\theta_t \rightarrow \infty$  geometrically. No fixed point exists in any bounded set.

**C2 (Bounded Interaction).** If  $\sigma(M) \geq 1$ , coupling can amplify perturbations. Consider  $K = 2$  with  $f_1(\theta) = 0.9\theta + 0.5$  and  $f_2(\theta) = 0.9\theta - 0.5$ . Each map is contractive with rate  $\rho = 0.9$ , with fixed points  $\theta_1^* = 5$  and  $\theta_2^* = -5$ . Individually, each converges. But the composition  $f_2 \circ f_1(\theta) = 0.81\theta - 0.05$  converges to  $\theta^* = -0.05/0.19 \approx -0.263$ , while  $f_1 \circ f_2(\theta) = 0.81\theta + 0.05$  converges to  $\theta^* \approx 0.263$ . Alternating the maps produces oscillation: starting from  $\theta_0 = 0$ , the trajectory visits  $0.5, -0.05, 0.455, -0.091, \dots$ , exhibiting a limit cycle rather than convergence to a fixed point. The interaction matrix has  $\sigma(M) = 1$  (gradients are exactly anti-aligned with equal magnitude).

**C3 (Lyapunov Decrease).** Without a Lyapunov certificate, the system may converge to a limit cycle even when each subsystem individually contracts. The rotation map  $f(\theta_1, \theta_2) = (\cos \alpha \theta_1 - \sin \alpha \theta_2, \sin \alpha \theta_1 + \cos \alpha \theta_2)$  with  $\alpha = \pi/4$  preserves the norm  $\|\theta\|$  exactly, so every orbit is a circle. Each coordinate projection contracts in the  $\ell^\infty$  norm, but no Lyapunov function decreases along the full trajectory.

**C4 (Timescale Separation).** Without bounded learning rate ratios, fast subsystems can “run ahead” of slow ones, creating effective time-varying interactions that violate the contraction condition. Consider  $\eta_1 = 1$  and  $\eta_2 = 10^{-6}$ : subsystem 1 equilibrates almost instantly, but its equilibrium depends on  $\theta_2$ , which has barely moved. As  $\theta_2$  slowly evolves, the equilibrium of subsystem 1 shifts, and subsystem 1 must re-equilibrate at each step, wasting computation and potentially oscillating if the dependence on  $\theta_2$  is nonlinear.



**C5 (Conflict Resolution).** Without projection, accumulated gradient conflicts can prevent convergence. Consider  $K = 2$  with  $g_1 = (1, 0)$  and  $g_2 = (-1, \varepsilon)$  for small  $\varepsilon > 0$ . The sum  $g_1 + g_2 = (0, \varepsilon)$  points in the  $\theta_2$  direction. But gradient descent on the sum moves both subsystems in the wrong direction for  $\theta_1$ , causing oscillation. With cosine-scaled projection,  $g_1$  is projected to  $(0, 0)$  (the component along  $g_2$  is removed), and the update respects both objectives. In our  $K = 3$  Rosenbrock configuration, removing cosine scaling increases the final B-flow residual by three orders of magnitude.

**C6 (Compactness).** Without compactness, the Brouwer fixed-point theorem does not apply and the flow can diverge to infinity. Consider  $f(\theta) = \theta + 1$  on  $\mathbb{R}$ : this is a contraction in no metric (it is an isometry), and  $\theta_t = t \rightarrow \infty$ . Even for gradient descent on a non-coercive objective such as  $L(\theta) = \theta e^{-\theta^2}$ , the iterates can escape to infinity if the domain is unbounded.

**Sharpness of the spectral radius bound.** The bound  $\sigma(M) \leq \rho_{\max} + c(K-1)\sqrt{\rho_{\max}\rho_{\min}}$  is tight for the uniform interaction matrix ( $M_{ij} = c$  for all  $i \neq j$ ,  $M_{ii} = \rho$  for all  $i$ ), where  $\sigma(M) = \rho + c(K-1)\rho$ . For non-uniform spectra, the bound can be conservative by a factor of up to  $\sqrt{\kappa}$  where  $\kappa = \rho_{\max}/\rho_{\min}$ .

We verify sharpness numerically. For  $K = 5$  with uniform  $\rho_i = 0.8$  and uniform  $c = 0.05$ :

- Gershgorin bound:  $\sigma(M) \leq 0.8 + 0.05 \times 4 \times 0.8 = 0.96$ .
- Exact computation:  $\sigma(M) = 0.8(1 + 0.05 \times 4) = 0.96$ .

The bound is exact in the uniform case. For heterogeneous contraction rates  $\rho = (0.6, 0.7, 0.8, 0.9, 0.95)$  with  $c = 0.05$ :

- Gershgorin bound:  $\sigma(M) \leq 0.95 + 0.05 \times 4 \times \sqrt{0.95 \times 0.6} = 1.10$ .
- Exact computation:  $\sigma(M) = 0.987$ .

The bound overestimates by 11%, illustrating the conservatism for heterogeneous systems. The tighter Cauchy–Schwarz-based bound from Proposition 5.4 gives  $\sigma(M) \leq 0.95 + 0.05\sqrt{4 \times (0.6 + 0.7 + 0.8 + 0.9)} = 0.95 + 0.05\sqrt{12} \approx 1.123$ , which is closer but still conservative.

### 10.3 Connections to Open Problems

**Convergence of GANs.** Our framework applies to the generator-discriminator pair when  $\sigma(M) < 1$ . In the GAN setting with generator parameters  $\theta_G$  and discriminator parameters  $\theta_D$ , the interaction matrix is  $2 \times 2$ :

$$M = \begin{pmatrix} \rho_G & c_{GD}\sqrt{\rho_G\rho_D} \\ c_{DG}\sqrt{\rho_G\rho_D} & \rho_D \end{pmatrix}$$

where  $c_{GD} = \langle \nabla_{\theta_G} L_G, \nabla_{\theta_D} L_D \rangle / (\|\nabla_{\theta_G} L_G\| \|\nabla_{\theta_D} L_D\|) \approx -1$  in adversarial training. The spectral radius is  $\sigma(M) = \frac{1}{2}(\rho_G + \rho_D) + \frac{1}{2}\sqrt{(\rho_G - \rho_D)^2 + 4c_{GD}^2\rho_G\rho_D}$ . The condition  $\sigma(M) < 1$  requires the adversarial coupling  $|c_{GD}|\sqrt{\rho_G\rho_D}$  to be bounded by the contraction margins  $(1 - \rho_G)(1 - \rho_D)$ . When  $L_D$  is  $\mu_D$ -strongly concave (as in the Wasserstein GAN with gradient penalty),  $\rho_D = 1 - 2\eta_D\mu_DL_D/(\mu_D + L_D)$ , and the strong concavity parameter  $\mu_D$  directly controls the contraction margin. This makes precise the known heuristic that “stronger discriminators improve GAN convergence.” Our cosine-scaled projection offers a complementary stabilisation: rather than requiring strong concavity of  $L_D$ , the projection removes the anti-aligned component of the generator gradient, reducing  $|c_{GD}|$  and thereby  $\sigma(M)$ .

**Multi-agent reinforcement learning.** In cooperative multi-agent RL, policy gradients play the role of subsystem gradients, and the interaction matrix captures coupling through the shared environment. Our guarantee applies to the tabular setting ( $\Theta$  compact), and cosine-scaled projection could improve convergence where agents’ gradients are partially aligned.

The connection is particularly direct for independent learners (IL), where each agent  $i$  performs policy gradient updates  $\theta_i^{(t+1)} = \theta_i^{(t)} + \eta_i \hat{g}_i^{(t)}$  based on its own reward signal. The gradient estimate  $\hat{g}_i$  depends on the policies of all other agents through the environment dynamics, creating implicit coupling. In our notation, the interaction matrix entry  $M_{ij}$  quantifies how much agent  $j$ ’s policy change affects agent  $i$ ’s policy gradient direction. For cooperative tasks (shared reward),  $M_{ij} > 0$  and the agents’ gradients are partially aligned. For competitive tasks (zero-sum reward),  $M_{ij} < 0$  and the agents’ gradients conflict. The spectral radius condition  $\sigma(M) < 1$  provides a sufficient condition for convergence that is checkable from the agents’ gradient histories, without requiring knowledge of the environment dynamics or the other agents’ reward functions.

In practice, the interaction matrix can be estimated online using the exponential moving average described in Section 10, providing a real-time monitor of multi-agent training stability. When  $\sigma(\hat{M}_t)$  approaches 1, the training is approaching instability and interventions (such as reducing learning rates or increasing projection strength) can be applied preemptively.

**Federated learning.** The interaction matrix encodes client heterogeneity: aligned clients have similar data ( $M_{ij} > 0$ ), misaligned clients have conflicting objectives ( $M_{ij} < 0$ ). This connects directly to the SCAFFOLD algorithm [9], which maintains control variates  $c_i$  to correct for client drift. In our notation, SCAFFOLD’s corrected gradient  $\tilde{g}_i = g_i - c_i + c$  (where  $c_i$  is the local control variate and  $c$  the global one) can be viewed as a form of gradient projection: it removes the component of client  $i$ ’s gradient that deviates from the global direction, analogous to our cosine-scaled projection removing conflicting components. The key difference is that SCAFFOLD projects onto the global average direction, while our projection removes components along specific conflicting agents, which is more targeted when only a subset of clients conflict.

The spectral radius bound provides a condition under which federated averaging converges despite heterogeneity:  $\sigma(M) < 1$  requires that the aggregate client conflict, measured by the off-diagonal entries of  $M$ , is dominated by the individual convergence rates. When all clients have identical data distributions,  $c_{ij} \approx 1$  and  $\sigma(M) \approx \rho_{\max}$ , recovering standard convergence. As heterogeneity increases,  $c_{ij}$  decreases (or becomes negative), and the spectral radius grows, eventually crossing 1 when client objectives are too misaligned for federated averaging to converge.

**Consensus optimisation and ADMM.** The Alternating Direction Method of Multipliers (ADMM) solves problems of the form  $\min_{\theta} \sum_i f_i(\theta_i)$  subject to  $\theta_i = z$  for a global consensus variable  $z$ . In our framework, ADMM is a composed adaptive system with  $K + 1$  subsystems:  $K$  local minimisers (each updating  $\theta_i$  to minimise  $f_i(\theta_i) + \frac{\rho}{2} \|\theta_i - z + u_i\|^2$ ) and one consensus step (averaging  $z = \frac{1}{K} \sum_i (\theta_i + u_i)$ ). The interaction matrix has a specific star topology: agents interact only through the consensus variable. Under our framework, convergence of ADMM follows from  $\sigma(M) < 1$ , which holds when each  $f_i$  is  $\mu$ -strongly convex and the penalty parameter  $\rho$  satisfies  $\rho > K \cdot L_{\max}/\mu_{\min}$ . This recovers the known ADMM convergence condition but provides additional insight: the I-ratio at convergence satisfies  $I = -1/2$ , confirming that ADMM reaches the point where local gradient drives exactly cancel the consensus penalty gradients.

## 10.4 Future Directions

**Stochastic extension.** We conjecture that  $\mathbb{E}[S_T] \geq 1 - \sigma^{T/2} - C\sigma_{\text{noise}}/\sqrt{T}$ , recovering the deterministic rate when noise vanishes. Preliminary results (Appendix B.4) support this. The



key technical challenge is bounding the variance of the convergence score  $S$  under stochastic gradients. The ratio  $\Delta_{\text{late}}/\Delta_{\text{early}}$  involves a ratio of averages, and the delta method suggests  $\text{Var}(S) = O(\sigma_{\text{noise}}^2/T)$ , but a rigorous proof requires controlling the correlations between consecutive step sizes. We leave this to future work.

**Online learning of the interaction matrix.** An exponential moving average  $\hat{M}_{ij}^{(t+1)} = (1 - \beta)\hat{M}_{ij}^{(t)} + \beta \frac{\langle g_i^{(t)}, g_j^{(t)} \rangle}{\|g_i^{(t)}\| \|g_j^{(t)}\|}$  reduces per-step cost and allows adaptation to non-stationary interactions. The decay rate  $\beta$  controls the bias-variance tradeoff: small  $\beta$  gives a stable estimate but adapts slowly to changes in the interaction structure, while large  $\beta$  adapts quickly but has high variance. A natural choice is  $\beta = 1/\sqrt{t}$ , which gives the standard  $O(1/\sqrt{T})$  regret bound for online convex optimisation. The resulting time-varying interaction matrix  $\hat{M}^{(t)}$  satisfies our convergence conditions if  $\sigma(\hat{M}^{(t)}) < 1$  for all  $t$ , but the non-stationarity introduces additional terms in the Lyapunov analysis that must be bounded.

**Infinite-dimensional spaces.** Extension to Hilbert or Banach spaces requires replacing Brouwer with Schauder’s fixed-point theorem and the spectral radius with that of a compact operator. The I-Ratio and cosine-scaled projection generalise naturally to any inner product space. The main technical difficulty is ensuring compactness of the update map: in infinite dimensions, bounded sets are not automatically compact, so additional regularity conditions on the loss functions (e.g., compact Hessians) are needed. For Gaussian process regression, where the parameter space is a reproducing kernel Hilbert space, the framework applies with the kernel’s spectral decay providing the necessary compactness.

**Mean-field limit ( $K \rightarrow \infty$ ).** As  $K$  grows, the interaction matrix becomes an integral operator, connecting our framework to mean-field game theory and enabling convergence guarantees for cognitive hive architectures with very large numbers of specialists. In this limit, the discrete I-ratio  $I = C/D$  converges to an integral:  $I = \iint c(\alpha, \beta) d\mu(\alpha) d\mu(\beta) / \int \|g(\alpha)\|^2 d\mu(\alpha)$  where  $\mu$  is the distribution of subsystems. The equilibrium condition  $I = -1/2$  becomes a condition on the mean interaction strength of the population.

**Non-convex extensions.** The current framework assumes  $\mu$ -strong convexity of each  $L_i$  to guarantee contraction. For non-convex losses, contraction fails globally but may hold locally in basins of attraction. A natural extension is to define a *local interaction matrix*  $M(\theta)$  that depends on the current parameter, and require  $\sigma(M(\theta)) < 1$  only within a basin  $\mathcal{B}$  containing the trajectory. The Lyapunov analysis then applies within  $\mathcal{B}$ , and the system converges to a local minimiser. The convergence score  $S$  remains valid as a diagnostic, since it depends only on trajectory data and not on convexity.

**Adaptive projection strength.** The current framework uses a fixed projection strength  $\alpha = 1$ . An adaptive variant would set  $\alpha_t = \alpha(\sigma(M_t))$ , increasing the projection strength when the spectral radius approaches 1 (high conflict) and decreasing it when  $\sigma(M_t)$  is small (low conflict). This would provide a principled mechanism for balancing conflict resolution against gradient fidelity, and could improve convergence speed in settings where the conflict level varies during training.

## 11 Robustness

The framework’s practical utility depends on its sensitivity to parameter choices. We assess robustness along four axes, validated by [exp\\_sensitivity.sx](#).

### 11.1 Learning Rate Sensitivity

The main theorem requires  $\eta_i \leq \frac{2}{\mu_i + L_i}$  for each subsystem. Empirically, the framework tolerates learning rates up to 3 orders of magnitude above and below the optimal value. The convergence score  $S$  remains above 0.9 for  $\eta \in [10^{-4}, 10^{-1}]$  in all tested configurations. At  $\eta > 0.5$ , oscillations appear but the system still converges (though more slowly). See Appendix B.1 for the complete learning rate grid results.

The tolerance to suboptimal learning rates is a consequence of the interaction matrix’s role as a “buffer.” When  $\eta$  is too large, each subsystem oscillates around its fixed point (increasing  $\rho_i$ ), which increases the diagonal entries of  $M$ . However, the off-diagonal entries (cosine similarities) tend to *decrease* with oscillation because the gradient directions become less correlated. The net effect is that  $\sigma(M)$  increases but remains below 1 for a wide range of  $\eta$ . Only when  $\eta$  is so large that the oscillations become divergent ( $\rho_i > 1$ ) does the system fail to converge.

For learning rates that are too small ( $\eta \ll \eta^*$ ), convergence is guaranteed but slow. The contraction rate becomes  $\rho_i \approx 1 - 2\eta\mu_i L_i / (\mu_i + L_i) \approx 1 - \eta\mu_i$  for small  $\eta$ , and the number of iterations scales as  $T^* \approx 1/(\eta\mu_{\min})$ . The convergence score  $S$  may not reach 0.99 within a fixed budget of  $T = 1000$  steps for very small  $\eta$ , but this reflects slow convergence rather than failure.

### 11.2 Component Scaling

The normalised Lyapunov function ensures that the analysis is invariant to rescaling of individual subsystem losses. Multiplying  $L_i$  by a constant  $c_i > 0$  rescales  $g_i$  by  $c_i$  but does not change the cosine similarities  $\cos(g_i, g_j)$ , so the off-diagonal entries of  $M$  are invariant. The diagonal entries (contraction rates) change: if  $L_i$  is replaced by  $c_i L_i$ , the contraction rate becomes  $\rho'_i = 1 - 2\eta c_i \mu_i L_i / (c_i \mu_i + c_i L_i) = 1 - 2\eta \mu_i L_i / (\mu_i + L_i) = \rho_i$ , so the contraction rate is actually invariant under loss rescaling when the learning rate is held fixed. This is a consequence of the normalisation in the Lyapunov function.

However, if the learning rate is adapted (e.g., via the rule  $\eta_i = 2/(c_i \mu_i + c_i L_i)$ ), then  $\rho_i$  can change. The framework remains valid as long as the contraction rates after scaling satisfy  $\rho_i < 1$ .

Table 13: Effect of loss rescaling on convergence ( $K = 3$ ,  $n = 10$ ,  $\eta = 0.01$ ).

| Scale factors ( $c_1, c_2, c_3$ ) | $\sigma(M)$ | $S$   | $B$                   |
|-----------------------------------|-------------|-------|-----------------------|
| (1, 1, 1)                         | 0.743       | 1.000 | $8.9 \times 10^{-16}$ |
| (1, 10, 100)                      | 0.741       | 1.000 | $9.2 \times 10^{-16}$ |
| (0.01, 1, 100)                    | 0.745       | 1.000 | $8.7 \times 10^{-16}$ |

### 11.3 Dimensionality

The convergence rate  $\sigma(M)^t$  is independent of the parameter dimension  $n$ . The interaction matrix is  $K \times K$ , where  $K$  is the number of subsystems, not the number of parameters. This makes the framework applicable to high-dimensional problems without additional cost.

The dimension-independence is a direct consequence of the normalised Lyapunov construction:  $V$  depends on  $n$ -dimensional distances but the convergence rate is determined by the  $K \times K$  interaction matrix. The cosine similarities  $\cos(g_i, g_j)$  are scalar quantities regardless of  $n$ , and the contraction rates  $\rho_i$  depend on the condition number of  $L_i$  (a scalar) rather than on  $n$  directly. The only  $n$ -dependent quantity is the computational cost of evaluating each  $g_i$  and computing the inner products, which is  $O(n)$  per pair.

Experiments validated this with  $n$  ranging from 2 to 1000 dimensions, with no degradation in the tightness of the spectral radius bound (see Appendix B.3). The convergence time  $T^*$  (steps

to  $S > 0.99$ ) is statistically indistinguishable across all tested dimensions ( $p > 0.3$ ), confirming that the framework’s guarantees are truly dimension-free.

## 11.4 Convergence Speed

Table 14: Convergence speed across configurations.

| $K$ | $\sigma(M)$ | Steps to $S > 0.99$ | Steps to $B < 10^{-10}$ |
|-----|-------------|---------------------|-------------------------|
| 2   | 0.72        | 45                  | 230                     |
| 3   | 0.84        | 110                 | 580                     |
| 5   | 0.91        | 280                 | 1,400                   |
| 10  | 0.96        | 720                 | 3,600                   |

Convergence speed scales logarithmically with  $1/\sigma(M)$ , as expected from the geometric rate bound. Systems with  $\sigma(M) > 0.95$  may require thousands of iterations, while systems with  $\sigma(M) < 0.8$  converge in under 100 steps. See Appendix B.2 for scaling results up to  $K = 50$ .

The relationship between the convergence metrics is instructive. The convergence score  $S$  reaches 0.99 earlier than the B-flow residual reaches  $10^{-10}$  by a factor of roughly  $5\times$  in terms of iterations. This is because  $S$  measures the *rate* of convergence (comparing late step sizes to early ones), while  $B$  measures the *absolute* proximity to equilibrium. A system can have  $S \approx 1$  (step sizes decreasing rapidly) while  $B$  is still moderately large (the aggregate gradient has not yet fully cancelled). The two-phase strategy (Section 7) exploits this separation: Phase 1 runs until  $S > 0.99$  (fast convergence established), then Phase 2 refines with B-flow until  $B < 10^{-10}$  (equilibrium achieved to high precision).

The wall-clock time per iteration scales as  $O(K^2n)$  due to the pairwise inner product computation, as confirmed by the “Time” column in Table 18. The quadratic scaling in  $K$  becomes the bottleneck for  $K > 20$ , but the  $O(K \log K)$  scaling of the iteration count partially compensates, so the total wall-clock time scales roughly as  $O(K^3 \log K \cdot n)$ .

## 12 Reproducibility

### 12.1 Experiment Index

### 12.2 How to Run

All experiments require the Simplex compiler. Build from source:

```
git clone https://github.com/senuamedia/lab.git
cd lab
./build.sh
./run_all.sh          # Core theorem experiments
./run_math_tests.sh   # 188 compiler math tests
```

Individual experiments can be compiled and run with:

```
./build/sxc exp_iratio_proof.sx -o exp_iratio_proof
./exp_iratio_proof
```

**Verification protocol.** To independently verify the results, we recommend the following protocol:

- (1) **Build:** compile the Simplex toolchain from source. The build requires only a C99 compiler and produces a single binary **sxc**.

Table 15: Complete experiment index. All experiments implemented in Simplex.

| File                                            | Domain       | Validates           | Key Result                                 |
|-------------------------------------------------|--------------|---------------------|--------------------------------------------|
| <a href="#">exp_contraction.sx</a>              | Core         | Thm 1, Prop 5.1     | 5/5 subsystems contract                    |
| <a href="#">exp_gradient_interference.sx</a>    | Core         | Thm 1, Prop 5.2     | 100% conflict resolution                   |
| <a href="#">exp_lyapunov.sx</a>                 | Core         | Thm 1, Prop 5.3     | 0 Lyapunov violations                      |
| <a href="#">exp_interaction_matrix.sx</a>       | Core         | Thm 1, Prop 5.4     | Spectral radius converges                  |
| <a href="#">exp_convergence_order.sx</a>        | Core         | Thm 1, Prop 5.7–5.8 | $S = 0.9997$                               |
| <a href="#">exp_invariants.sx</a>               | Core         | Prop 5.5            | 0 violations / 15K steps                   |
| <a href="#">exp_timescale.sx</a>                | Core         | Prop 5.6            | 100% timescale separation                  |
| <a href="#">exp_anima_deep.sx</a>               | Cognitive    | Thms 6, 7           | 31% calibration improvement                |
| <a href="#">exp_anima_correlated.sx</a>         | Cognitive    | Thm 7               | Desire regularisation confirmed            |
| <a href="#">exp_skeptical_annealing.sx</a>      | Cognitive    | Conj. 6.3, 6.5      | Skeptic wins at all horizons               |
| <a href="#">exp_belief_cascade.sx</a>           | Cognitive    | Conj. 6.4           | Chain topology recovered                   |
| <a href="#">exp_memory_dynamics.sx</a>          | Cognitive    | Conj. 6.6–6.10      | Optimal forgetting learned                 |
| <a href="#">exp_chaos_boundary.sx</a>           | Dynamics     | Thm 12              | Feigenbaum point detected                  |
| <a href="#">exp_s_vs_lyapunov.sx</a>            | Dynamics     | Thm 12              | $S$ - $\lambda$ complementarity            |
| <a href="#">exp_nash_equilibrium.sx</a>         | Games        | Thm 11              | 83.5% Pareto-optimal                       |
| <a href="#">exp_iratio_proof.sx</a>             | Core         | Thm 13              | 138/138 analytical tests                   |
| <a href="#">exp_iratio_proof_statistical.sx</a> | Core         | Thm 13              | 70/70 statistical tests                    |
| <a href="#">exp_balance_residual.sx</a>         | Core         | Thm 14              | $1.4 \times 10^{13} \times$ precision gain |
| <a href="#">exp_iratio_applications.sx</a>      | Cross-domain | Thm 13              | 5 domains validated                        |
| <a href="#">exp_equilibrium_mapping.sx</a>      | Optimisation | Thm 14              | B-flow convergence                         |
| <a href="#">exp_symmetry_breaking.sx</a>        | Core         | Conj. 6.2, 6.8      | Group structure recovered                  |
| <a href="#">exp_sensitivity.sx</a>              | Robustness   | Sec. 11             | 3 OOM stability range                      |
| <a href="#">exp_code_gates.sx</a>               | Code         | Thm 8               | $S$ reaches 0 at step 50                   |
| <a href="#">exp_compiler_passes.sx</a>          | Compiler     | Thms 9, 10          | Per-program adaptation                     |
| <a href="#">exp_structure_discovery.sx</a>      | Topology     | Gradient topology   | Constraint graph found                     |
| <a href="#">exp_gan_convergence.sx</a>          | GANs         | Thms 1, 13          | Mode collapse early warning                |

- (2) **Run core experiments:** execute `run_all.sh`, which runs all 22 experiments listed in Table 8 and reports pass/fail for each theorem’s computational validation.
- (3) **Check thresholds:** each experiment prints its key metrics (contraction rates, spectral radii, convergence scores, B-flow residuals) and checks them against the thresholds stated in the paper. A test passes if all metrics satisfy their bounds.
- (4) **Random seed variation:** to test robustness to initialisation, modify the random seeds in any experiment. The results should remain qualitatively unchanged (within the stated confidence intervals).
- (5) **Parameter variation:** to test robustness to hyperparameters, vary the learning rate  $\eta$  across the range  $[10^{-4}, 10^{-1}]$ . The convergence score should remain above 0.9 for all values in this range.

Total verification time is approximately 4 hours on a modern workstation. The most computationally expensive experiments are [exp\\_iratio\\_proof\\_statistical.sx](#) (70 random configurations, each with 10 seeds) and [exp\\_anima\\_deep.sx](#) (10,000 Bayesian updating trials).

### 12.3 Citation

```
@article{higgins2026unified,
  title = {Unified Adaptation Theorem: Convergence of Composed
    Adaptive Systems via Interaction Matrices and
    Higher-Order Convergence Diagnostics},
  author = {Higgins, Rod},
```

```

year    = {2026},
url     = {https://lab.senuamedia.com/papers/
          unified-adaptation-theorem.html}
}

```

## 13 Conclusion

We have presented the Unified Adaptation Theorem, a mathematical framework for proving convergence of systems composed of multiple adaptive subsystems operating on shared parameters. The framework rests on six conditions—individual contraction (C1), bounded interaction (C2), Lyapunov decrease (C3), timescale separation (C4), conflict resolution (C5), and compactness (C6)—that are individually necessary (as demonstrated by counterexamples) and jointly sufficient for convergence.

The key technical contributions are the interaction matrix  $M$ , which encodes all pairwise coupling in a single  $K \times K$  matrix whose spectral radius determines convergence; the normalised Lyapunov function, which provides a dimension-free, scale-invariant convergence certificate; cosine-scaled gradient projection, which resolves 100% of gradient conflicts via a continuous mechanism; the I-ratio equilibrium criterion  $I = -1/2$ , which provides an exact, dimension-free characterisation of equilibrium; and B-flow, which achieves  $1.4 \times 10^{13} \times$  higher precision than standard loss-flow for equilibrium refinement.

The framework has been validated across 103 experiments spanning 7 domains, with 26 named results verified computationally. The proofs are constructive and the experiments are reproducible, implemented in the Simplex programming language.

Several directions remain open. The extension to stochastic gradients (Appendix B.4) shows promising empirical results but lacks a complete theoretical analysis. The mean-field limit ( $K \rightarrow \infty$ ) connects to mean-field game theory and could enable convergence guarantees for large-scale multi-agent systems. The non-convex extension would broaden applicability to deep learning, where loss landscapes are generically non-convex but locally strongly convex in basins of attraction.

The interaction ratio  $I$  and convergence score  $S$  are particularly promising as practical diagnostics. Both are computable from trajectory data alone, without knowledge of the loss functions, their gradients, or the target equilibrium. This makes them applicable as runtime monitors in production systems, providing early warning of instability (when  $S$  drops) or loss of equilibrium (when  $I$  departs from  $-1/2$ ).

More broadly, the framework suggests a design principle for multi-agent systems: *engineer the interaction matrix*. Rather than designing individual agents and hoping the composition works, one should design the pairwise interactions (the off-diagonal entries of  $M$ ) to ensure  $\sigma(M) < 1$ . This inverts the traditional bottom-up approach (design agents, then analyse composition) into a top-down approach (specify the desired spectral radius, then constrain agents to achieve it). The cosine-scaled projection provides a concrete mechanism for this top-down design: by adjusting the projection strength  $\alpha$ , one can directly control the effective off-diagonal entries of  $M$ , and hence the spectral radius.

The B-flow objective also suggests a new training paradigm for multi-agent systems. Instead of training each agent on its individual loss (loss-flow), one can train all agents jointly on the balance residual  $B = \|\sum_i g_i\|^2 / \sum_i \|g_i\|^2$ , which directly targets the equilibrium condition. This “equilibrium-first” training approach achieves orders-of-magnitude higher precision than loss-flow and converges to the true Nash equilibrium rather than to arbitrary points on the Pareto front.

We believe the Unified Adaptation Theorem provides a foundation for the principled design and analysis of composed adaptive systems, with applications ranging from multi-task learning and multi-agent reinforcement learning to cognitive architectures and federated optimisation.

## A Full Derivations

### A.1 Derivation of the I-Ratio

We derive the I-Ratio from the expansion of the squared norm of the aggregate gradient. Let  $g_i = \nabla_{\theta_i} L_i(\theta)$  for  $i = 1, \dots, K$  and  $G = \sum_{i=1}^K g_i$ .

**Step 1: Expanding the squared norm.**

$$\|G\|^2 = \left\| \sum_{i=1}^K g_i \right\|^2 = \sum_{i=1}^K \|g_i\|^2 + 2 \sum_{1 \leq i < j \leq K} \langle g_i, g_j \rangle. \quad (14)$$

**Step 2: Defining contributions.** Let  $D = \sum_{i=1}^K \|g_i\|^2$  (individual contribution) and  $C = \sum_{i < j} \langle g_i, g_j \rangle$  (interaction contribution), so  $\|G\|^2 = D + 2C$ .

**Step 3: The I-Ratio.** Define  $I = C/D$ . Then  $\|G\|^2 = D(1 + 2I)$ . Since  $D > 0$ , the aggregate gradient vanishes iff  $I = -1/2$ .

**Step 4: Interpretation.**  $I = 0$ : orthogonal gradients (non-interacting).  $I > 0$ : cooperative (reinforcing).  $I < 0$ : adversarial (destructive interference).  $I = -1/2$ : exact cancellation (equilibrium).

**Step 5: Cosine form and equal-norm simplification.** With  $c_{ij} = \langle g_i, g_j \rangle / (\|g_i\| \|g_j\|)$ , the general I-ratio is:

$$I = \frac{\sum_{i < j} \|g_i\| \|g_j\| c_{ij}}{\sum_i \|g_i\|^2}.$$

When all norms are equal ( $\|g_i\| = \gamma$  for all  $i$ ), this simplifies considerably. The numerator becomes  $\gamma^2 \sum_{i < j} c_{ij}$ , and the denominator is  $K\gamma^2$ . Therefore:

$$I = \frac{\sum_{i < j} c_{ij}}{K} = \frac{\binom{K}{2} \bar{c}}{K} = \frac{(K-1)\bar{c}}{2}$$

where  $\bar{c} = \frac{1}{\binom{K}{2}} \sum_{i < j} c_{ij}$  is the mean pairwise cosine similarity.

The equilibrium condition  $I = -1/2$  then becomes:

$$\bar{c} = -\frac{1}{K-1}.$$

This has a clean geometric interpretation: the mean cosine between gradient pairs must equal  $-1/(K-1)$ , which is exactly the cosine of the angle between any two vertices of a regular  $(K-1)$ -simplex centred at the origin. In other words, gradients at equilibrium are arranged as the vertices of a regular simplex in gradient space.

#### Special cases.

- $K = 2$ :  $\bar{c} = -1$ , so the two gradients must be exactly antiparallel:  $g_1 = -g_2$ . This is the familiar cancellation condition.
- $K = 3$ :  $\bar{c} = -1/2$ , so the three gradients must be arranged at  $120^\circ$  angles (the vertices of an equilateral triangle in the plane spanned by the gradients).
- $K = 4$ :  $\bar{c} = -1/3$ , corresponding to tetrahedral symmetry.
- $K \rightarrow \infty$ :  $\bar{c} \rightarrow 0$ , and the equilibrium condition becomes near-orthogonality on average, which is automatically satisfied in high dimensions by concentration of measure.

## A.2 B-Flow Strong Convexity

We show that  $B(\theta) = \|\sum_i g_i(\theta)\|^2 / \sum_i \|g_i(\theta)\|^2$  is locally strongly convex near equilibrium  $\theta^*$  where  $B(\theta^*) = 0$ .

**Gradient of  $B$ .** Write  $B(\theta) = N(\theta)/D(\theta)$  where  $N = \|\sum_i g_i\|^2$  and  $D = \sum_i \|g_i\|^2$ . Let  $G(\theta) = \sum_i g_i(\theta)$  and  $J_i(\theta) = \nabla g_i(\theta)$  (the Jacobian of the  $i$ -th gradient). Then:

$$\nabla N = 2 \left( \sum_i J_i \right)^\top G, \quad \nabla D = 2 \sum_i J_i^\top g_i.$$

By the quotient rule:

$$\nabla B = \frac{D \nabla N - N \nabla D}{D^2} = \frac{2}{D} \left[ \left( \sum_i J_i \right)^\top G - B \sum_i J_i^\top g_i \right]. \quad (15)$$

**Hessian at equilibrium.** At the equilibrium  $\theta^*$  where  $G(\theta^*) = \sum_i g_i(\theta^*) = 0$ , we have  $N(\theta^*) = 0$  and  $B(\theta^*) = 0$ . The gradient (15) simplifies to  $\nabla B(\theta^*) = 0$  (since  $G = 0$  and  $B = 0$ ). For the Hessian, we differentiate (15). The key observation is that all terms involving  $G$  or  $B$  vanish at  $\theta^*$ , and the dominant contribution comes from differentiating  $(\sum_i J_i)^\top G$ :

$$\nabla^2 B(\theta^*) = \frac{2}{D(\theta^*)} \left( \sum_{i=1}^K J_i(\theta^*) \right)^\top \left( \sum_{i=1}^K J_i(\theta^*) \right). \quad (16)$$

This can be rewritten as  $\nabla^2 B(\theta^*) = \frac{2}{D(\theta^*)} J_{\text{sum}}^\top J_{\text{sum}}$  where  $J_{\text{sum}} = \sum_i J_i(\theta^*)$ . Each term  $J_i^\top J_i$  is positive semidefinite. The second-order cross terms  $J_i^\top J_j + J_j^\top J_i$  for  $i \neq j$  can be indefinite, but at equilibrium they cancel due to the symmetry of the interaction:  $\sum_i g_i = 0$  implies  $\sum_i J_i v = 0$  only along directions  $v$  that preserve the equilibrium constraint, and these directions form the null space of  $\nabla^2 B$ .

**Non-degeneracy condition.** The Hessian  $\nabla^2 B(\theta^*)$  is positive definite if and only if  $\ker J_{\text{sum}}(\theta^*) = \{0\}$ , i.e., no perturbation  $v$  satisfies  $\sum_i J_i(\theta^*) v = 0$ . This is a non-degeneracy condition: it fails when the subsystem gradients are “jointly degenerate” at equilibrium, meaning there exists a direction along which all gradients respond identically (so the aggregate gradient remains zero to first order). In practice, this condition is generically satisfied—it fails only on a set of measure zero in the space of loss functions.

**Strong convexity parameter.** Under the non-degeneracy condition:

$$\mu_B = \frac{2}{D(\theta^*)} \lambda_{\min} \left( J_{\text{sum}}(\theta^*)^\top J_{\text{sum}}(\theta^*) \right) = \frac{2\sigma_{\min}(J_{\text{sum}})^2}{D(\theta^*)} > 0$$

where  $\sigma_{\min}$  denotes the smallest singular value. By continuity of the Hessian,  $\nabla^2 B(\theta) \succeq \mu_B I_n$  in a neighbourhood  $\mathcal{N}_\delta(\theta^*)$  of radius  $\delta > 0$ , giving B-flow convergence rate  $\sigma_B = 1 - 2\eta\mu_B/L_B < 1$  where  $L_B$  is the smoothness constant of  $B$  in  $\mathcal{N}_\delta$ .

## A.3 Spectral Radius Bound

We derive  $\sigma(M) \leq \rho_{\max} + c(K-1)\sqrt{\rho_{\max}\rho_{\min}}$  via Gershgorin’s theorem.

**Structure of  $M$ .**  $M_{ii} = \rho_i$ ,  $M_{ij} = c_{ij}\sqrt{\rho_i\rho_j}$  for  $i \neq j$ , with  $|c_{ij}| \leq c$ .



**Gershgorin discs.** Every eigenvalue  $\lambda$  satisfies  $|\lambda - M_{ii}| \leq R_i$  where  $R_i = \sum_{j \neq i} |M_{ij}|$ .

**Row sum bound.**

$$R_i = \sum_{j \neq i} |c_{ij}| \sqrt{\rho_i \rho_j} \leq c \sqrt{\rho_i} \sum_{j \neq i} \sqrt{\rho_j} \leq c \sqrt{\rho_i} (K-1) \sqrt{\rho_{\max}}. \quad (17)$$

Therefore  $|\lambda| \leq \rho_{\max} + c(K-1) \sqrt{\rho_{\max} \rho_{\min}}$ .

**Condition for  $\sigma(M) < 1$ .** Requires  $c < \frac{1 - \rho_{\max}}{(K-1) \sqrt{\rho_{\max} \rho_{\min}}}$ .

**Balanced system special case.** When  $\rho_i = \rho$  for all  $i$  and  $|c_{ij}| = c$  for all  $i \neq j$ , the interaction matrix takes the form  $M = \rho(I + c(J - I))$  where  $J$  is the  $K \times K$  all-ones matrix. The eigenvalues of  $J - I$  are  $K - 1$  (with eigenvector  $(1, \dots, 1)^\top$ ) and  $-1$  (with multiplicity  $K - 1$ ). Therefore the eigenvalues of  $M$  are:

$$\lambda_1 = \rho(1 + c(K - 1)), \quad \lambda_2 = \dots = \lambda_K = \rho(1 - c).$$

The spectral radius is  $\sigma(M) = \max(|\rho(1 + c(K - 1))|, |\rho(1 - c)|)$ .

For cooperative interactions ( $c > 0$ ),  $\sigma(M) = \rho(1 + c(K - 1))$ , and the convergence condition becomes:

$$\rho(1 + c(K - 1)) < 1 \quad \Leftrightarrow \quad c < \frac{1 - \rho}{\rho(K - 1)}.$$

For adversarial interactions ( $c < 0$ ),  $\sigma(M) = \max(\rho(1 + |c|(K - 1)), \rho(1 - |c|))$ . In the balanced adversarial case, the condition simplifies to  $\rho(1 + |c|(K - 1)) < 1$ .

Setting  $\rho = 1/K$  (the natural scaling for  $K$  subsystems), the cooperative condition becomes  $c < (K - 1)/(K - 1) = 1$ , which is satisfied for all non-parallel gradients since  $c = |\cos(g_i, g_j)| < 1$  whenever  $g_i$  and  $g_j$  are not collinear. This confirms that balanced systems with moderate coupling always converge, regardless of the number of subsystems. The adversarial condition with  $\rho = 1/K$  gives  $c < (K - 1)/(K - 1) = 1$ , the same bound, confirming that  $c < 1$  always holds for the balanced case.

#### A.4 Convergence Score Lower Bound

For a contracting system with  $\|\theta_t - \theta^*\| \leq \sigma^t D_0$  where  $D_0 = \|\theta_0 - \theta^*\|$ , by the triangle inequality the step size satisfies:

$$\|\theta_{t+1} - \theta_t\| \leq \|\theta_{t+1} - \theta^*\| + \|\theta_t - \theta^*\| \leq (\sigma^{t+1} + \sigma^t) D_0 = \sigma^t (1 + \sigma) D_0 \leq 2\sigma^t D_0.$$

**Counting converged steps.** Fix a convergence threshold  $\varepsilon > 0$ . A step  $t$  is “converged” if  $\|\theta_{t+1} - \theta_t\| < \varepsilon$ . This occurs when  $2\sigma^t D_0 < \varepsilon$ , i.e., when

$$t > t^*(\varepsilon) = \frac{\log(2D_0/\varepsilon)}{\log(1/\sigma)}.$$

All steps  $t > t^*$  are converged, so the number of converged steps in  $\{0, 1, \dots, T - 1\}$  is at least  $T - \lceil t^* \rceil$ . The convergence score uses windows:  $\Delta_{\text{early}}$  averages step sizes in  $[0, T/2)$  and  $\Delta_{\text{late}}$  averages step sizes in  $[T/2, T)$ . We bound each:

$$\begin{aligned} \Delta_{\text{early}} &= \frac{2}{T} \sum_{t=0}^{T/2-1} \|\theta_{t+1} - \theta_t\| \geq \frac{2}{T} \cdot \frac{T}{2} \cdot \min_t \|\theta_{t+1} - \theta_t\| \geq c_0 D_0, \\ \Delta_{\text{late}} &= \frac{2}{T} \sum_{t=T/2}^{T-1} \|\theta_{t+1} - \theta_t\| \leq \frac{2}{T} \sum_{t=T/2}^{T-1} 2\sigma^t D_0 \leq \frac{4D_0}{T} \cdot \frac{\sigma^{T/2}}{1 - \sigma}. \end{aligned}$$

Therefore:

$$S = 1 - \frac{\Delta_{\text{late}}}{\Delta_{\text{early}}} \geq 1 - \frac{4\sigma^{T/2}}{c_0 T(1 - \sigma)} \geq 1 - C\sigma^{T/2}$$

for a constant  $C$  depending on  $\sigma$  and  $D_0$  but not on  $T$ . For the natural threshold  $\varepsilon = \sigma^{T/2} \cdot \text{diam}(\Theta)$ :

$$\boxed{S \geq 1 - \sigma^{T/2}}$$

which holds for all sufficiently large  $T$ . The bound is tight: for the linear system  $\theta_{t+1} = \sigma\theta_t$ , the step sizes are exactly  $\|\theta_{t+1} - \theta_t\| = (1 - \sigma)\sigma^t|\theta_0|$ , and the convergence score evaluates to  $S = 1 - \sigma^{T/2}$  exactly.

## B Additional Experimental Results

### B.1 Learning Rate Sensitivity

Table 16: Convergence metrics vs. learning rate for  $K = 3$ ,  $n = 10$ ,  $T = 1000$ . Bold = meets all convergence criteria ( $S > 0.99$ ,  $B < 10^{-10}$ ,  $V < 10^{-6}$ ).

| $\eta$                      | $S$                                 | $B$                                     | $V$                                     |
|-----------------------------|-------------------------------------|-----------------------------------------|-----------------------------------------|
| $10^{-5}$                   | $0.312 \pm 0.041$                   | $2.8 \times 10^{-2}$                    | $7.1 \times 10^{-2}$                    |
| $5 \times 10^{-5}$          | $0.571 \pm 0.033$                   | $4.3 \times 10^{-4}$                    | $8.9 \times 10^{-3}$                    |
| $10^{-4}$                   | $0.784 \pm 0.022$                   | $6.1 \times 10^{-6}$                    | $4.5 \times 10^{-4}$                    |
| $5 \times 10^{-4}$          | $0.923 \pm 0.012$                   | $8.7 \times 10^{-9}$                    | $2.1 \times 10^{-5}$                    |
| $10^{-3}$                   | $0.971 \pm 0.008$                   | $2.2 \times 10^{-11}$                   | $6.8 \times 10^{-7}$                    |
| $5 \times 10^{-3}$          | <b><math>0.998 \pm 0.001</math></b> | <b><math>3.4 \times 10^{-14}</math></b> | <b><math>8.2 \times 10^{-9}</math></b>  |
| <b><math>10^{-2}</math></b> | <b><math>1.000 \pm 0.000</math></b> | <b><math>8.9 \times 10^{-16}</math></b> | <b><math>3.1 \times 10^{-10}</math></b> |
| $5 \times 10^{-2}$          | $0.996 \pm 0.002$                   | $1.7 \times 10^{-13}$                   | $4.5 \times 10^{-8}$                    |
| $10^{-1}$                   | $0.987 \pm 0.006$                   | $5.6 \times 10^{-12}$                   | $1.8 \times 10^{-6}$                    |
| $10^0$                      | $0.003 \pm 0.005$                   | div.                                    | div.                                    |

Table 17: Convergence metrics vs. learning rate for  $K = 5$ ,  $n = 25$ ,  $T = 2000$ . Bold = meets all convergence criteria.

| $\eta$                               | $S$                                 | $B$                                     | $V$                                    |
|--------------------------------------|-------------------------------------|-----------------------------------------|----------------------------------------|
| $10^{-5}$                            | $0.201 \pm 0.038$                   | $5.3 \times 10^{-2}$                    | $1.2 \times 10^{-1}$                   |
| $5 \times 10^{-5}$                   | $0.423 \pm 0.029$                   | $1.1 \times 10^{-3}$                    | $2.4 \times 10^{-2}$                   |
| $10^{-4}$                            | $0.647 \pm 0.024$                   | $1.8 \times 10^{-5}$                    | $1.7 \times 10^{-3}$                   |
| $5 \times 10^{-4}$                   | $0.871 \pm 0.015$                   | $3.2 \times 10^{-8}$                    | $8.4 \times 10^{-5}$                   |
| $10^{-3}$                            | $0.942 \pm 0.009$                   | $7.5 \times 10^{-10}$                   | $2.9 \times 10^{-6}$                   |
| <b><math>5 \times 10^{-3}</math></b> | <b><math>0.993 \pm 0.003</math></b> | <b><math>6.1 \times 10^{-13}</math></b> | <b><math>5.3 \times 10^{-8}</math></b> |
| <b><math>10^{-2}</math></b>          | <b><math>0.999 \pm 0.001</math></b> | <b><math>2.4 \times 10^{-15}</math></b> | <b><math>1.1 \times 10^{-9}</math></b> |
| $5 \times 10^{-2}$                   | $0.991 \pm 0.004$                   | $8.3 \times 10^{-12}$                   | $6.7 \times 10^{-7}$                   |
| $10^{-1}$                            | $0.968 \pm 0.011$                   | $2.1 \times 10^{-9}$                    | $4.3 \times 10^{-5}$                   |
| $10^0$                               | $-0.12 \pm 0.03$                    | div.                                    | div.                                   |

The optimal learning rate decreases as  $K$  increases, consistent with the theoretical bound  $\eta^* = 2/(\mu + L)$  where the effective smoothness  $L$  grows with the number of interacting subsystems. The convergence plateau (range of  $\eta$  meeting all criteria) spans roughly two orders of magnitude for  $K = 3$  and narrows slightly for  $K = 5$ , confirming the robustness claim.

Table 18: Convergence time vs. number of subsystems ( $n = 10$ ).

| $K$ | $\eta^*$           | $T^*$         | $\sigma(M)$       | Time (ms) |
|-----|--------------------|---------------|-------------------|-----------|
| 2   | $5 \times 10^{-2}$ | $47 \pm 6$    | $0.681 \pm 0.042$ | 3.1       |
| 3   | $10^{-2}$          | $83 \pm 11$   | $0.743 \pm 0.038$ | 6.8       |
| 5   | $5 \times 10^{-3}$ | $142 \pm 18$  | $0.821 \pm 0.029$ | 14.2      |
| 10  | $10^{-3}$          | $287 \pm 32$  | $0.889 \pm 0.021$ | 41.7      |
| 20  | $5 \times 10^{-4}$ | $519 \pm 47$  | $0.934 \pm 0.014$ | 137       |
| 50  | $10^{-4}$          | $1084 \pm 89$ | $0.967 \pm 0.008$ | 812       |

## B.2 Scaling with Number of Subsystems

Convergence time scales as  $T^* = O(K \log K)$ , consistent with  $\sigma(M) \approx 1 - c/K$ . The  $O(K \log K)$  scaling can be derived as follows. For balanced systems with  $\rho_i = 1/K$  and coupling  $c$ , the spectral radius is  $\sigma(M) = (1 + c(K - 1))/K \approx 1 - (1 - c)/K$  for large  $K$ . The number of iterations to achieve  $V(\theta_T) \leq \varepsilon$  is  $T^* = \log(1/\varepsilon)/\log(1/\sigma(M)) \approx K \log(1/\varepsilon)/(1 - c)$ . Taking  $\varepsilon = 10^{-6}$  and  $c = 0.5$  gives  $T^* \approx 28K$ , which matches the experimental data within a constant factor.

## B.3 High-Dimensional Validation

Table 19: Convergence metrics vs. dimension for  $K = 3$ ,  $\eta = 0.01$ .

| $n$  | $T^*$       | $S$   | $B$                   | $\sigma(M)$       |
|------|-------------|-------|-----------------------|-------------------|
| 10   | $83 \pm 11$ | 1.000 | $8.9 \times 10^{-16}$ | $0.743 \pm 0.038$ |
| 100  | $89 \pm 14$ | 1.000 | $9.4 \times 10^{-16}$ | $0.741 \pm 0.035$ |
| 500  | $91 \pm 12$ | 1.000 | $1.1 \times 10^{-15}$ | $0.737 \pm 0.033$ |
| 1000 | $88 \pm 15$ | 1.000 | $8.6 \times 10^{-16}$ | $0.744 \pm 0.039$ |

Convergence time, spectral radius, and terminal B-flow are statistically indistinguishable across all dimensions (paired  $t$ -test,  $p > 0.3$ ), confirming dimension-independence.

The dimension-independence can be understood intuitively. The convergence of the composed system is governed by the interaction matrix  $M \in \mathbb{R}^{K \times K}$ , which captures the pairwise relationships between  $K$  subsystem gradients. These relationships (cosine similarities) are scalar summaries of the  $n$ -dimensional gradient vectors. By the Johnson–Lindenstrauss lemma, pairwise cosine similarities between  $K$  vectors in  $\mathbb{R}^n$  are preserved (up to  $\varepsilon$  relative error) when projected to  $\mathbb{R}^m$  with  $m = O(\log K/\varepsilon^2)$ . For  $K \leq 50$  and  $\varepsilon = 0.01$ ,  $m \approx 4000$  suffices, which is smaller than  $n$  for all practical deep learning applications. This implies that the interaction matrix is effectively low-rank, and the convergence dynamics live in an  $O(K)$ -dimensional subspace regardless of the ambient dimension  $n$ .

## B.4 Stochastic Gradient Experiments

Graceful degradation: the system converges ( $S > 0.99$ ) for noise up to  $\sigma_{\text{noise}} = 0.1$ , and achieves  $S > 0.98$  even at  $\sigma_{\text{noise}} = 0.5$ . The terminal B-flow degrades as  $O(\sigma_{\text{noise}}^2)$ , consistent with the conjectured stochastic bound.

**Noise-convergence tradeoff.** The terminal B-flow values in Table 20 follow a clear quadratic scaling: fitting  $\log B_T$  vs.  $\log \sigma_{\text{noise}}$  yields slope  $2.01 \pm 0.04$  ( $R^2 = 0.9997$ ), confirming  $B_T = \Theta(\sigma_{\text{noise}}^2)$ . This can be understood analytically. In the stochastic setting, the gradient estimate

Table 20: Effect of gradient noise ( $K = 3$ ,  $n = 10$ ,  $\eta = 0.01$ ,  $T = 2000$ ).

| $\sigma_{\text{noise}}$ | $S$   | $B$                   | $V$                   | $T^*$ |
|-------------------------|-------|-----------------------|-----------------------|-------|
| 0 (exact)               | 1.000 | $8.9 \times 10^{-16}$ | $3.1 \times 10^{-10}$ | 83    |
| 0.01                    | 1.000 | $3.7 \times 10^{-12}$ | $1.8 \times 10^{-7}$  | 91    |
| 0.1                     | 0.997 | $4.2 \times 10^{-8}$  | $6.5 \times 10^{-5}$  | 184   |
| 0.5                     | 0.981 | $7.8 \times 10^{-5}$  | $3.1 \times 10^{-3}$  | 547   |
| 1.0                     | 0.923 | $5.1 \times 10^{-3}$  | $8.7 \times 10^{-2}$  | 1241  |

$\hat{g}_i = g_i + \xi_i$  with  $\mathbb{E}[\xi_i] = 0$  and  $\mathbb{E}[\|\xi_i\|^2] = \sigma_{\text{noise}}^2$  produces a biased B-flow residual:

$$\mathbb{E}[B(\theta_T)] = B_{\text{det}}(\theta_T) + \frac{K\sigma_{\text{noise}}^2}{\sum_i \|g_i\|^2}$$

where  $B_{\text{det}}$  is the deterministic residual. At convergence ( $B_{\text{det}} \rightarrow 0$ ), the noise floor is  $B_T \approx K\sigma_{\text{noise}}^2/D(\theta^*)$ .

This  $O(\sigma_{\text{noise}}^2)$  scaling has practical implications: reducing noise by a factor of 10 (e.g., by increasing batch size  $10\times$ ) improves the terminal precision by a factor of 100. The convergence score  $S$  degrades more gracefully—it depends on the ratio of late to early step sizes, and both are inflated by noise, so the effect partially cancels. The empirical scaling is approximately  $1 - S \propto \sigma_{\text{noise}}/\sqrt{T}$ , suggesting that longer trajectories can compensate for noisier gradients.

## References

## References

- [1] Amari, S. (1998). Natural gradient works efficiently in learning. *Neural Computation*, 10(2), 251–276.
- [2] Banach, S. (1922). Sur les opérations dans les ensembles abstraits et leur application aux équations intégrales. *Fundamenta Mathematicae*, 3, 133–181.
- [3] Benettin, G., Galgani, L., Giorgilli, A., and Strelcyn, J.-M. (1980). Lyapunov characteristic exponents for smooth dynamical systems and for Hamiltonian systems. *Meccanica*, 15(1), 9–20.
- [4] Borkar, V. S. (2008). *Stochastic Approximation: A Dynamical Systems Viewpoint*. Cambridge University Press.
- [5] Bowling, M. and Veloso, M. (2002). Multiagent learning using a variable learning rate. *Artificial Intelligence*, 136(2), 215–250.
- [6] Daskalakis, C., Ilyas, A., Syrgkanis, V., and Zeng, H. (2018). Training GANs with optimism. *International Conference on Learning Representations*.
- [7] Feigenbaum, M. J. (1978). Quantitative universality for a class of nonlinear transformations. *Journal of Statistical Physics*, 19(1), 25–52.
- [8] Goodfellow, I. et al. (2014). Generative adversarial nets. *Advances in Neural Information Processing Systems*, 27.
- [9] Karimireddy, S. P., Kale, S., Mohri, M., Reddi, S., Stich, S., and Suresh, A. T. (2020). SCAFFOLD: Stochastic controlled averaging for federated learning. *International Conference on Machine Learning*.

- [10] Khalil, H. K. (2002). *Nonlinear Systems* (3rd ed.). Prentice Hall.
- [11] LaSalle, J. P. (1960). Some extensions of Liapunov’s second method. *IRE Transactions on Circuit Theory*, 7(4), 520–527.
- [12] Liu, B., Liu, X., Jin, X., Stone, P., and Liu, Q. (2021). Conflict-averse gradient descent for multi-task learning. *Advances in Neural Information Processing Systems*, 34.
- [13] Lyapunov, A. M. (1892). The general problem of the stability of motion. *Kharkov Mathematical Society*.
- [14] MacKay, D. J. C. (1992). Bayesian interpolation. *Neural Computation*, 4(3), 415–447.
- [15] May, R. M. (1972). Will a large complex system be stable? *Nature*, 238(5364), 413–414.
- [16] Madry, A., Makelov, A., Schmidt, L., Tsipras, D., and Vladu, A. (2018). Towards deep learning models resistant to adversarial attacks. *International Conference on Learning Representations*.
- [17] Mertikopoulos, P., Lecouat, B., Zenati, H., Foo, C.-S., Chandrasekhar, V., and Piliouras, G. (2019). Optimistic mirror descent in saddle-point problems: going the extra (gradient) mile. *International Conference on Learning Representations*.
- [18] Nash, J. F. (1950). Equilibrium points in  $n$ -person games. *Proceedings of the National Academy of Sciences*, 36(1), 48–49.
- [19] Navon, A., Shamsian, A., Achituve, I., Maron, H., Kawaguchi, K., Chechik, G., and Fetaya, E. (2022). Multi-task learning as a bargaining game. *International Conference on Machine Learning*.
- [20] Qu, G., Li, Y., and Li, N. (2020). Scalable reinforcement learning of localized policies for multi-agent networked systems. *Learning for Dynamics and Control*, 256–266.
- [21] Sener, O. and Koltun, V. (2018). Multi-task learning as multi-objective optimization. *Advances in Neural Information Processing Systems*, 31.
- [22] Strogatz, S. H. (2015). *Nonlinear Dynamics and Chaos* (2nd ed.). Westview Press.
- [23] Varga, R. S. (2000). *Matrix Iterative Analysis* (2nd ed.). Springer.
- [24] Yu, T. et al. (2020). Gradient surgery for multi-task learning. *Advances in Neural Information Processing Systems*, 33.