

Article

Not peer-reviewed version

Beyond the Data Paradigm: Freedom Intelligence and the Structural Laws of Navigability

[Gonçalo Melo de Magalhães](#) *

Posted Date: 5 March 2026

doi: 10.20944/preprints202603.0429.v1

Keywords: data-centric AI; model-centric AI; navigability; freedom intelligence; structural learning; perception-distortion; gradient law; freedom intelligence training; world models; causal learning; physics-informed AI



Preprints.org is a free multidisciplinary platform providing preprint service that is dedicated to making early versions of research outputs permanently available and citable. Preprints posted at Preprints.org appear in Web of Science, Crossref, Google Scholar, Scilit, Europe PMC.

Copyright: This open access article is published under a [Creative Commons CC BY 4.0 license](#), which permit the free download, distribution, and reuse, provided that the author and preprint are cited in any reuse.

Disclaimer/Publisher's Note: The statements, opinions, and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions, or products referred to in the content.

Article

Beyond the Data Paradigm: Freedom Intelligence and the Structural Laws of Navigability

Gonalo Melo de Magalhães

Independent Researcher; Porto, Portugal; hi@planta.design

Abstract

Machine learning's dominant paradigm—whether model-centric or data-centric—treats intelligence as the extraction of statistical patterns from behavioral records. This approach has delivered remarkable engineering feats. Yet something foundational is missing. Data is not reality: it is a finite record of trajectories through reality. A photograph of a river is not the river's law. This paper argues that the data paradigm conflates measurement with mechanism, capturing where systems have been rather than why they go there. We propose an alternative grounded in the Architecture of Freedom Intelligence (AFI), which identifies navigability—the structural availability of paths—as the primary organizing principle of all complex systems. The Law of Freedom, $F = P/D$, states that navigational capacity equals differentiation capacity (Perception, P) divided by structural resistance (Distortion, D). Under this framework, intelligence is not pattern memorization but distortion navigation: all systems move according to $dx/dt = -P(x) \cdot \nabla D(x)$, following gradients of resistance scaled by perceptual capacity. We demonstrate that this gradient law is structurally identical to Fick's diffusion, Berg–Brown chemotaxis, Ohm's law, and gradient descent—revealing a deep structural unity that the data paradigm treats as coincidental analogy. Nature does not train on labeled datasets: ants, neurons, immune cells, and ecological populations navigate through calibrated heuristics on Perception and Distortion fields, not through backpropagation over historical trajectories. This observation motivates a fundamental reconceptualization of what training should accomplish. We propose Freedom Intelligence Training (FIT): a learning paradigm oriented toward learning P and D fields directly, rather than fitting statistical correlations over behavioral snapshots. FIT rests on five predictions: (i) models trained on P – D fields require exponentially less data than pattern-extraction models; (ii) generalization improves because P – D fields encode causal structure; (iii) out-of-distribution performance improves because navigability laws transfer across domains; (iv) interpretability is natural since every prediction decomposes into ΔP and ΔD contributions; (v) the exploration–exploitation transition is quantifiable as the coefficient of variation of the Freedom field crossing 1.0. We provide ten falsification criteria and position FIT within the emerging landscape of world models, physics-informed learning, and causal inference. This is a theoretical proposal; a complete experimental roadmap is provided.

Keywords: data-centric AI; model-centric AI; navigability; freedom intelligence; structural learning; perception-distortion; gradient law; freedom intelligence training; world models; causal learning; physics-informed AI

1. Introduction: The Two Crises of the Data Paradigm

The history of machine learning is, in large part, a history of debates about data. The first great debate—model-centric vs. data-centric AI—asked whether intelligence emerges from better algorithms or better datasets. The emerging consensus is that both matter, and that the two approaches are complementary rather than competing [1,2]. Yet this consensus, while practically useful, conceals a deeper question that neither tradition has adequately addressed: What is data, structurally, and what can it—and cannot—tell us about the systems that generate it?

Current AI development faces two simultaneous crises that the model-centric/data-centric debate cannot resolve. The first is the scaling crisis: large language models require exponentially more compute and data to achieve linear performance improvements, and multiple analyses suggest this trajectory is approaching physical and economic limits [3]. The second is the generalization crisis: models trained on massive datasets routinely fail on distributions they have not seen, hallucinate physical impossibilities, and cannot reason causally about the world they describe in fluent language [4,5]. These two crises share a common root: the data paradigm treats statistical patterns as proxies for structural laws, when in fact they are only imperfect shadows of those laws.

This paper proposes a different starting point. We argue that reality is not primarily a dataset—it is a navigability field. Systems, whether physical, biological, or computational, do not traverse reality by consulting stored trajectories. They navigate by responding to structural gradients: moving away from resistance, toward availability, guided by their differentiation capacity. A river does not remember where other rivers went. An ant colony does not store labeled examples of good versus bad paths. A neuron does not train on a supervised classification task. These systems obey a structural law—what we term the Law of Freedom—that governs navigability directly, without reference to historical data.

The Law of Freedom, proposed as the central equation of the Architecture of Freedom Intelligence (AFI) framework [6], states:

$$F = P/D \tag{1}$$

where F is Freedom (structural navigability), P is Perception (differentiation capacity in bit-cycles), and D is Distortion (structural resistance to traversal). Systems move according to the gradient law:

$$dx/dt = -P(x) \cdot \nabla D(x) \tag{2}$$

This law is not a metaphor. It is structurally identical—under appropriate mappings—to Fick's diffusion, Berg–Brown chemotaxis, Ohm's law, and gradient descent [7–9]. The fact that these five independently derived equations share the same mathematical form is not coincidence: it is evidence of an underlying structural law governing all navigating systems.

If this structural law exists—if $F = P/D$ governs navigability across domains—then the goal of intelligent learning is not to memorize behavioral trajectories, but to calibrate the Perception and Distortion fields that generate those trajectories. This shift in goal implies a shift in what we mean by training, what we count as a model, and what we measure as success. We call this shift Freedom Intelligence Training (FIT).

This paper proceeds as follows. Section 2 performs a structural analysis of what data is and is not. Section 3 introduces the Law of Freedom and its uniqueness proof. Section 4 shows how the gradient law unifies five canonical equations. Section 5 examines how natural systems learn without datasets. Section 6 develops Freedom Intelligence Training as a formal alternative. Section 7 positions FIT within related paradigms. Section 8 applies FIT to the five original AFI theses. Section 9 provides falsification criteria. Section 10 addresses objections. Section 11 discusses limits and the experimental roadmap. Section 12 concludes.

2. What Data Is—And Is Not: A Structural Analysis

Before proposing an alternative to the data paradigm, we must be precise about what data is. This requires distinguishing three levels of description that are routinely conflated in machine learning practice.

2.1. The Three Levels

Level 1: Reality. Reality consists of systems in motion through possibility spaces. A bird migrates, a market adjusts, a protein folds, a pedestrian navigates a corridor. These motions are governed by structural laws—physical forces, gradient fields, resistance landscapes. Reality at this level is dynamic, continuous, and causally structured.

Level 2: Flow. Flow is the spatiotemporal trajectory of a system through its possibility space. Flow has direction (the gradient law determines it), speed (proportional to P and inversely proportional to D), and curvature (determined by the shape of the D field). Flow is reality in motion, not a record of reality.

Level 3: Data. Data is a finite sample of flow: a discrete set of (state, timestamp) pairs recorded by an observer. Data is always a projection of flow onto a measurement instrument. It inherits the structural biases of that instrument (sensor resolution, sampling rate, observation window), the biases of the flow that generated it (which paths were traversed), and the biases of the reality that structured the flow (which paths existed).

The data paradigm equates Level 3 (data) with Level 1 (reality), treating statistical patterns in behavioral records as if they were the laws governing those behaviors. This is the fundamental error. The correct inversion is: data reveals Level 3, which reflects Level 2, which is governed by Level 1.

2.2. The River Photograph Problem

Consider a river. Photographed at 100 different locations over five years, you will have a large dataset of river-states. A model trained on this dataset can predict, with some accuracy, what the river will look like at location x on date y . This is data-centric AI applied to rivers.

But this model cannot tell you why the river flows where it does. It cannot tell you what would happen if you removed a dam, added a tributary, or changed the rainfall pattern. It does not know about topography, hydraulic gradients, or fluid dynamics. It knows only where this particular river has been.

The river's actual law is $dx/dt = -\nabla V(x)$, where V is the hydraulic potential field (a Distortion field). The data paradigm learns a statistical approximation of this law from historical trajectories. Freedom Intelligence learns the law directly: what is $V(x)$, how does it shape gradients, and therefore where will the river go? A model knowing $V(x)$ can recompute $F = P/D$ under a modified landscape—dam removal, changed rainfall—and accurately predict the new flow. A behavioral-pattern model cannot.

2.3. Data as Compressed Trajectories

Let S be a state space, $D(x)$ a distortion field on S , and $P(x)$ a perception field. The flow is governed by $dx/dt = -P(x) \cdot \nabla D(x)$. A dataset is a finite collection of pairs $\{(x_i, t_i)\}$ sampled from this flow—a lossy compression of the flow, which is itself a consequence of the structural fields P and D .

This characterization has three immediate implications. First, training on data is, at best, an indirect route to learning P and D . Second, the amount of data needed to recover P and D via indirect trajectory sampling scales with the complexity of the D field and the dimensionality of S —explaining why deep learning requires massive datasets: it is recovering a high-dimensional structural field from sparse trajectory samples. Third, systems with simple D fields can be learned from few trajectories; systems with complex D fields require either massive data or direct D -field learning.

The Freedom Intelligence answer to the data crisis is therefore not more data, nor better data—it is different data: samples not of behavioral trajectories, but of the resistance and perception fields that generate those trajectories.

3. The Law of Freedom: $F = P/D$

3.1. Formal Statement

The Law of Freedom states that the structural navigability of a system equals its differentiation capacity divided by its structural resistance. We define:

Perception (P): The capacity of a system to distinguish among alternative continuation paths. Formally: $P = \log_2(N) \times T$, where N is the number of distinguishable states the system can differentiate

(sensor resolution, available options), and T is the temporal depth of the system's predictive horizon (autocorrelation decay, lookahead steps). Units: bit-cycles. Minimum value: $P = 1$ (one state, one step).

Distortion (D): The structural resistance to traversal. In single-factor domains, D may be measured as time ($D = t_{\text{traverse}} / t_{\text{baseline}}$), energy ($D = E_{\text{required}} / E_{\text{baseline}}$), or error rate ($D = 1/(1-p_{\text{error}})$). In complex multi-factor systems: $D = R^{\alpha} \cdot O^{\beta} \cdot T_{\text{burden}}^{\gamma} \cdot C^{\delta} \cdot M^{\epsilon}$, where R is redundancy, O is obligation (checkpoints), T_{burden} is temporal load, C is complexity, and M is inertia. Each factor is dimensionless (normalized to baseline = 1). The multiplicative composition is essential: independent barriers compound—if $D_A = 3$ and $D_B = 2$, combined $D = 6$, not 5; any single infinite barrier blocks the system regardless of how easy other barriers are; and log-space becomes additive, enabling standard log-linear regression [10].

Freedom (F): The resulting structural navigability: $F = P / D$.

3.2. Uniqueness Proof

The ratio form $F = P/D$ is not arbitrary. Under three axioms, it is the unique form satisfying all three simultaneously:

(C1) Monotonicity: F increases in P , decreases in D . More differentiation capacity and less resistance both increase navigability.

(C2) Scale-covariance: $F(\lambda P, \lambda D) = F(P, D)$. Navigability is unchanged when both P and D are scaled by the same factor.

(C3) Dimensional separability: $F(P, D) = g(P) \cdot h(D)$. The Perception and Distortion contributions factorize—they are independently measurable and independently modifiable.

Under these axioms, the unique functional form is $F = c \cdot (P/D)^{\alpha}$. Setting $\alpha = 1$ as the minimal hypothesis: $F = P/D$. The additive form $F = P - D$ is excluded because it permits negative Freedom (incoherent). The multiplicative form $F = P \times D$ is excluded because it implies navigability increases with resistance (absurd). The uniqueness makes the law falsifiable: if empirical data require $\alpha \neq 1$, the law generalizes to $F = (P/D)^{\alpha}$ without changing its structure.

3.3. Three Navigational Regimes

The law takes different forms depending on a system's Perception capacity:

Passive Regime ($P = 1$): $F = 1/D$. The system does not differentiate among paths. Freedom varies inversely with resistance alone. Examples: water through a membrane (Fick's diffusion), current through a resistor (Ohm's law), a ball rolling downhill. This is the Principle of Least Action: when differentiation is absent, the system follows the steepest descent of the Distortion field.

Active Regime ($P > 1$, P independent of D): $F = P/D$. The system differentiates but does not adapt its perception to the resistance landscape. Doubling P doubles F ; doubling D halves F .

Intelligent Regime (P responds to D): $F_i = (P_{\text{ext}} \times P_{\text{rec}}) / (D_{\text{ext}} \times D_{\text{int}})$. The system recursively adapts its perception. P_{ext} is external perception (sensor-derived), P_{rec} is recursive perception (internal modeling), D_{ext} is environmental resistance, D_{int} is internal distortion (computational cost of recursive analysis). D_{ext} and D_{int} compose multiplicatively: independent barrier classes compound. The Intelligence Paradox emerges here: increasing recursive analysis ($\uparrow P_{\text{rec}}$) increases the numerator but simultaneously increases D_{int} . Net benefit requires the marginal return on perception to exceed the marginal cost in distortion.

4. The Gradient Law and the Deep Unity of Navigation

4.1. The Gradient Law

The scalar law $F = P/D$ specifies how much navigability a system possesses at a given state, but not which direction it moves. Direction emerges when Distortion varies spatially. Let $D(x)$ be the distortion field at state x . The gradient $\nabla D(x)$ points toward increasing resistance. Systems move opposite the gradient, scaled by Perception: $dx/dt = -P(x) \cdot \nabla D(x)$. The negative sign encodes

preference for lower resistance. Multiplication by $P(x)$ scales responsiveness: high-Perception systems respond to shallow gradients; low-Perception systems respond only to steep ones.

4.2. Five Structural Identities

A theoretical law claiming to unify navigational dynamics must pass a stringent test: structural identity—not mere analogy—with the canonical equations governing those domains. We demonstrate five such identities.

Table 1. Structural identities between the gradient law and five canonical equations.

Domain/Equation	P maps to	D maps to	Regime
Fick's diffusion: $J = -D_{\text{diff}} \cdot \nabla c$	D_{diff} (diffusion coefficient)	$-1/c$ (inverse concentration)	Passive
Berg–Brown chemotaxis: biased random walk	Receptor sensitivity	Inverse concentration nutrient	Active
Gradient descent: $\theta_{t+1} = \theta_t - \eta \nabla L$	η (learning rate)	L (loss function)	Passive/Active
Ohm's law: $I = -(1/R) \cdot \nabla V$	$1/R$ (conductance)	V (voltage)	Passive
Langevin (low noise): $dx = -\nabla U \cdot dt$	1	U (potential energy)	Passive

These are not metaphors: the mathematical forms are identical under the stated mappings. Fick's law obtains by setting $D(x) = -1/c(x)$ and $P = D_{\text{diff}}$. Ohm's law obtains by setting $D = V$ (voltage) and $P = 1/R$ (conductance). Gradient descent obtains by setting $P = \eta$ (the learning rate) and $D = L$ (the loss landscape) [9,11]. This identity is the formal content of the claim that $F = P/D$ is not domain-specific but structural.

This structural unity has a direct implication for machine learning: gradient descent—the core of virtually all deep learning—is itself a navigational system operating in the Passive Regime of Freedom Intelligence. The learning rate η is Perception; the loss landscape L is Distortion; the optimization trajectory is the system's path through the field. Freedom Intelligence proposes that we should also learn the shape of this Distortion field directly, rather than relying solely on gradient traversal.

4.3. The Counterfactual Test

A critical test of the gradient law is its counterfactual prediction: in regions where $\nabla D = 0$ (uniform Distortion), there should be no directional bias regardless of P . The system should exhibit random walk statistics. This has been confirmed across domains: ant colonies in environments with uniform path lengths show no directional preference; particle swarm optimizers in flat landscapes exhibit random drift; gradient descent at zero-gradient critical points produces no parameter update. This directional symmetry-breaking is a unique prediction of the gradient law that purely statistical models do not make.

5. Nature's Learning: Heuristics Without Datasets

5.1. The Biological Contrast

Machine learning achieves intelligence by consuming labeled examples. A convolutional neural network for image classification may train on millions of annotated images. A large language model trains on hundreds of billions of tokens, effectively consuming the recorded output of human civilization. This is the data paradigm operating at maximum scale.

Nature, by contrast, produces navigating systems of extraordinary sophistication without labeled datasets, backpropagation, or internet-scale training corpora. A foraging ant learns to prefer shorter paths over longer ones through the stigmergic accumulation and evaporation of pheromone—a physical process that encodes the Distortion field of the environment into the Perception field of the colony. An immune cell navigates toward a pathogen through chemotaxis: sensing a chemical gradient (the Distortion field) and moving along it (the gradient law). A bird navigating by magnetoreception calibrates its Perception of a physical field that encodes structural information about the Distortion landscape.

These are not crude approximations of deep learning. They are, in many cases, more sample-efficient, more robust, more generalizable, and more energy-efficient than current AI systems. An ant colony solves Traveling Salesman instances in near-real time with hundreds of agents and microseconds of computation per agent—comparable efficiency to specialized combinatorial solvers at a fraction of the energy cost [12].

5.2. Three Natural Learning Strategies

Nature employs three distinct strategies for learning the P-D landscape, all of which differ fundamentally from statistical pattern extraction.

Strategy 1: Stigmergic Field Encoding. Stigmergic systems (ant colonies, biofilms, slime molds) encode Distortion information in the physical environment itself. Pheromone trails are not data records: they are P-field updates. The conversion rate is governed by $dP_{\text{stig}}/dt = \kappa \cdot \tau(x,t) / D(x,t)$, where κ is conversion efficiency and τ is the cumulative trace. The information encoded is about Distortion structure (path lengths, obstacles), not behavioral history. The system learns the field, not the trajectory.

Strategy 2: Chemotactic Calibration. Chemotactic systems (bacteria, immune cells, growth cones) calibrate their Perception of chemical gradients through receptor adaptation. The Berg–Brown model of *E. coli* chemotaxis involves run-and-tumble dynamics where the tumbling rate decreases when the organism senses increasing nutrient concentration (decreasing D) [7]. The organism calibrates the sensitivity and temporal integration parameters of its Perception system to match the local gradient structure of the Distortion field. This is precisely the FIT objective: calibrate P to match D .

Strategy 3: Heuristic Field Compression. Many cognitive systems employ fast-and-frugal heuristics [13] that are structural rules for navigating the P-D ratio rather than statistical models of past behavior. Take-the-best, satisficing, and recognition-based decision making operate in bounded-Perception regimes by identifying high- D alternatives and eliminating them first. These heuristics do not require large training sets because they operate on the structure of the problem space, not on the history of traversals through it.

The common thread across all three strategies: nature learns about the field (P and D), not about the trajectories that the field generates. This is the natural motivation for Freedom Intelligence Training.

6. Freedom Intelligence Training: A Proposed Paradigm

6.1. Core Thesis

Freedom Intelligence Training (FIT) proposes the following shift in training objective:

Current paradigm: Given a dataset $\{(x_i, y_i)\}$, find parameters θ that minimize $\sum L(f_{\theta}(x_i), y_i)$. Train a model that maps inputs to outputs by fitting statistical correlations in behavioral records.

FIT paradigm: Given a system, learn the Perception field $P(x)$ and Distortion field $D(x)$ that govern its navigation. Predict behavior by applying the gradient law $dx/dt = -P(x) \cdot \nabla D(x)$ rather than by interpolating historical trajectories.

This is not a rejection of statistical learning: it is a redirection. Statistical methods remain essential for estimating $P(x)$ and $D(x)$ from observations. What changes is what is being estimated: not $p(y|x)$

(conditional behavioral distribution), but $P(x)$ (Perception field) and $D(x)$ (Distortion field). This is a shift in target, not a rejection of method.

6.2. FIT Architecture

A FIT model consists of three components:

Component 1—Perception Encoder (P-net): Maps the system's current state x to a differentiation capacity estimate $P(x)$. P-net learns how many alternatives does the system perceive at state x , and how far ahead can it effectively predict. Output: $P(x)$ as a scalar field over S .

Component 2—Distortion Estimator (D-net): Maps the system's current state x to a structural resistance estimate $D(x)$. D-net learns the shape of the Distortion landscape. Output: $D(x)$ as a scalar field over S , and $\nabla D(x)$ as a vector field.

Component 3—Navigation Engine: Given $P(x)$ and $D(x)$, applies the gradient law $dx/dt = -P(x) \cdot \nabla D(x)$ to predict the system's trajectory. For intelligent systems, applies the full formulation $F_i = (P_{\text{ext}} \times P_{\text{rec}}) / (D_{\text{ext}} \times D_{\text{int}})$.

The key insight: P-net and D-net are trained on different signals. P-net is trained on the system's discriminability—experiments or observations that reveal how many alternatives the system can distinguish, and at what temporal depth. D-net is trained on resistance measurements—time, energy, error, or composite burden across traversed paths. These are fundamentally different measurement protocols than behavioral labeling.

6.3. FIT Training Signals

FIT requires the following training signals, which are qualitatively different from labeled behavioral examples:

For P estimation: forced-choice experiments (how many alternatives can the system reliably distinguish?), temporal autocorrelation analysis (how far ahead does the system's state predict future states?), and sensor resolution measurements (what is the system's minimum discriminable difference?). These reveal the structure of the Perception field.

For D estimation: path cost measurements across a representative sample of state pairs (time, energy, error rate), resistance field reconstruction from observed deceleration or failure events, and comparative analysis of traversed vs. avoided paths (avoided paths typically have higher D). These reveal the structure of the Distortion field.

Neither type of training signal requires behavioral labels in the supervised learning sense. Both can often be obtained from physical measurements, engineering specifications, or first-principles analysis—reducing or eliminating dependence on large labeled datasets.

6.4. The Five FIT Predictions

FIT makes the following testable predictions:

Prediction 1—Data Efficiency: FIT models require exponentially less behavioral data than pattern-extraction models to achieve equivalent navigation performance. Justification: learning P and D fields is a lower-dimensional problem than learning a behavioral distribution over a high-dimensional state space. Falsifiable: compare FIT vs. behavioral-pattern model performance as a function of training set size.

Prediction 2—Causal Generalization: FIT models generalize better to structural interventions (changes to D) because they represent the mechanism, not the correlation. If D is modified (a barrier removed, a path added), the FIT model computes new $F = P/D$ over the modified landscape and correctly predicts new trajectories. A behavioral-pattern model will fail because the new trajectories were never in its training distribution [14]. Falsifiable: compare generalization after structural interventions.

Prediction 3—Out-of-Distribution Robustness: Because navigability laws transfer across domains (Fick, Berg–Brown, Ohm, and gradient descent all obey the same gradient law), a FIT model

calibrated on one domain should exhibit better out-of-distribution performance on structurally similar domains than a domain-specific behavioral model [15]. Falsifiable: cross-domain transfer experiments.

Prediction 4—Natural Interpretability: Every FIT prediction decomposes into a ΔP contribution (perception change) and a ΔD contribution (distortion change). This provides causal attribution by construction: we can answer not just 'what will happen' but 'why will it happen—because P increased, because D decreased, or both.' No post-hoc interpretability method is required. Falsifiable: compare explanation quality and fidelity against SHAP/LIME baselines.

Prediction 5—Quantifiable Exploration–Exploitation Transition: The coefficient of variation of the Freedom field $E(t) = \text{Var}[F(x,t)] / \text{Mean}[F(x,t)]$ provides a universal, algorithm-agnostic metric for the exploration–exploitation transition. When $E(t) > 1$, the system is in exploratory mode; when $E(t) < 1$, in exploitative mode; $E(t) = 1$ is the transition point. This prediction applies equally to ant colony optimization, gradient descent, biological foraging, and market dynamics. Falsifiable: track $E(t)$ across diverse navigating systems and verify that convergence events correlate with $E(t)$ crossing 1.0.

7. Positioning FIT within the AI Research Landscape

7.1. FIT vs. Model-Centric and Data-Centric AI

Model-centric AI optimizes the algorithm given fixed data; data-centric AI optimizes the data given a fixed algorithm [1,2]. Both operate within the assumption that the goal of AI is statistical pattern extraction from behavioral records. FIT challenges this assumption at the level of the training objective: the goal is not to extract patterns from behavioral records, but to learn the structural fields (P and D) that generate those records. This does not make FIT incompatible with either paradigm—statistical methods remain the primary tools for P-D estimation—but it shifts the target from $p(y|x)$ to $P(x)$ and $D(x)$.

7.2. FIT and World Models

The emerging world-model paradigm (LeCun's JEPA architecture, NVIDIA's Cosmos) attempts to move AI from pattern prediction to predictive world representation [16,17]. FIT is a specific formal instance of a world model: the world is represented as a P-D field pair, and the dynamics are governed by the gradient law. This gives FIT two advantages over general world-model architectures: a specific functional form ($F = P/D$) with a uniqueness proof, and a direct connection to five canonical physical equations, providing strong structural priors that reduce the learning problem's effective dimensionality.

7.3. FIT and Physics-Informed Machine Learning

Physics-Informed Machine Learning (PIML) integrates physical laws (typically PDEs) as constraints or inductive biases in neural network training [11]. FIT is structurally similar: the gradient law $dx/dt = -P(x) \cdot \nabla D(x)$ is a PDE that constrains the learned fields. However, FIT differs from PIML in scope: PIML typically applies domain-specific physical equations, while FIT applies a domain-agnostic structural law claimed to hold across all navigating systems. If this claim is correct, FIT provides a universal PIML-style constraint applicable anywhere the gradient law holds.

7.4. FIT and Causal AI

Causal AI (following Pearl's do-calculus and structural causal models) distinguishes interventional distributions $P(y|\text{do}(x))$ from observational distributions $P(y|x)$ [18]. A model that learns causal structure can answer 'what would happen if we intervened to change x?' FIT is a form of causal AI: the Distortion field $D(x)$ represents the causal mechanism governing system navigation, and $P(x)$ represents the causal capacity to perceive and respond to that mechanism. Changing D (a structural intervention) allows the FIT model to predict new navigational behavior through the same

gradient law. However, FIT makes a stronger claim than generic causal AI: it proposes that the specific functional form $F = P/D$ and the specific gradient law $dx/dt = -P \cdot \nabla D$ are universal across all navigating systems—a claim subject to cross-domain empirical test [15].

8. Freedom Intelligence, Training, and the Five AFI Theses

The AFI framework is built on five empirical theses. Each has direct implications for what FIT means and predicts.

8.1. Thesis 1: Freedom as Cause ($F \equiv F$)

Zero paths means no world. Freedom—the structural availability of paths—is the irreducible first condition for any coherent system. Before spacetime, matter, or law, the possibility of transition must exist [19,20]. Implication for FIT: data quality is not primarily a function of statistical richness, but of whether it samples the true structure of path availability in the domain. Data collected from systems operating in constrained path spaces may fail to reveal the true D landscape, no matter how much is collected.

8.2. Thesis 2: The Law of Freedom ($F = P/D$)

Navigability equals differentiation capacity divided by resistance. Implication for training: the fundamental quantities to estimate are P and D, not behavioral distributions. The R^2 of $F = P/D$ against navigation metrics (traversal speed, path quality, completion rate) across three or more independently measured domains is a pre-registered empirical threshold: $R^2 \geq 0.80$. If this threshold fails, FIT's theoretical foundation is falsified.

8.3. Thesis 3: Freedom as Return (The FLRP Architecture)

Coherent worlds require the generative order Freedom → Logic → Relations → Physics (FLRP). The Physics we observe is generated last, not first. Implication for training: the data paradigm trains on Physics (observable patterns) and attempts to infer Logic, Relations, and Freedom from the bottom up. FIT proposes the reverse: learn the Freedom structure (P-D fields) and derive Physics as a consequence. This predicts that FIT models will exhibit better extrapolation beyond observed physical regularities.

8.4. Thesis 4: Mutual Dependency

No FLRP component functions in isolation. Implication for training: a FIT model that learns D without P will fail (Passive Regime—it will have no directional sensitivity). A FIT model that learns P without D will fail (undefined navigation—perception without resistance has no gradient). Both fields must be estimated jointly, with their ratio ($F = P/D$) as the organizing principle.

8.5. Thesis 5: Space as Maximum Distortion

Physical space is the least adaptive substrate: the domain where Distortion is hardest to modify. Design intervention works by increasing P or decreasing effective D within fixed spatial constraints. Implication for training: for smart environment applications (building automation, industrial digital twins, urban mobility [21]), FIT models should focus on P-D estimation within fixed spatial D_{ext} , with interventions acting on P (signage, sensors, information) and D_{int} (procedural optimization).

9. Falsification Criteria

A theory without falsification criteria is not science. The following ten criteria specify the exact conditions under which FIT and its AFI foundation must be abandoned or fundamentally revised.

Table 2. Ten falsification criteria for AFI and FIT.

ID	Criterion	Falsification Condition
F1	Path irreducibility	A coherent transition system with zero available paths is demonstrated
F2	P/D proportionality	$R^2 < 0.80$ between measured P/D and navigation metrics in ≥ 3 independent domains
F3	Passive reduction	Physical systems with $P = 1$ deviate systematically from $F = 1/D$
F4	Gradient direction	Navigating systems move toward ∇D (increasing resistance) systematically
F5	Stigmergic conversion	Pheromone accumulation does not measurably increase effective Perception
F6	Cross-domain recurrence	$F = P/D$ proportionality holds in fewer than 5 independent domains
F7	FLRP ordering	Stable physical regularities consistently appear without prior path availability
F8	FIT efficiency	FIT models require as much or more data as behavioral models for equivalent task performance
F9	Causal generalization	FIT models fail to generalize better than behavioral models after structural interventions (changes to D)
F10	E(t) transition	The coefficient of variation of the Freedom field does not predict exploration-exploitation transitions across algorithm families

Criteria F1–F7 target the AFI framework. Criteria F8–F10 are specific to Freedom Intelligence Training. The framework collapses if F2 or F6 fail (the core proportionality law), or if both F8 and F9 fail (FIT has no advantage over behavioral learning).

10. Addressing the Strongest Objections

10.1. FIT is Just Physics-Informed ML Rebranded

This is the most technically informed objection. PIML applies domain-specific physical equations as constraints—Navier–Stokes for fluids, Maxwell for electromagnetics—each applying to one domain [11]. FIT claims a single equation (the gradient law) applies universally. The five structural identities in Section 4.2 are the evidence. If this structural identity is real—not metaphorical—then FIT is a specific hypothesis about the mathematical structure of the universe: all navigating systems obey the same gradient law. This is a testable empirical claim that PIML's framework-agnosticism does not make.

10.2. P and D Are Too Vague to Measure

This is the strongest operational objection. The response requires domain-specific operationalization with independent measurement protocols. For graph navigation: $P = \log_2(N_edges) \times T_horizon$ (measurable from graph structure and agent lookahead), $D = \text{path cost}$ (measurable from edge weights). The independence test: P and D must be modifiable independently. If increasing N_edges always co-varies with decreasing D, P and D are not independently measurable and the theory is not falsifiable in that domain. This independence test must be pre-registered before each validation experiment.

10.3. Deep Learning Already Learns Structural Features

Modern deep networks do learn features that generalize beyond surface statistics—convolutional networks learn edge detectors, transformers learn syntactic structures. This is correct, and FIT does not deny it. The claim is that (i) behavioral learning does so inefficiently, requiring massive datasets; (ii) the learned representations conflate P and D (they encode the ratio F, not the numerator and denominator independently); and (iii) this conflation prevents the natural interpretability and interventional generalization that FIT's explicit P-D decomposition enables [15,22]. The empirical question is whether explicit P-D learning achieves equivalent or better performance with less data.

10.4. Multiplicative D Has Too Many Free Parameters

The composite Distortion formula $D = R^{\alpha} \cdot O^{\beta} \cdot T^{\gamma} \cdot C^{\delta} \cdot M^{\epsilon}$ has five exponents. In many domains, fewer factors may dominate. Sobol sensitivity analysis should be applied to identify which factors carry the most first-order and total-order variance. Empirical evidence from building energy simulation suggests that temporal load (T_burden) and complexity (C) typically dominate, reducing the effective model to $D \approx T^{\gamma} \cdot C^{\delta}$. Furthermore, both the multiplicative and additive forms should be tested and compared. If the additive form achieves equivalent R^2 in a specific domain, this should be reported honestly—it does not falsify the theory but informs domain-specific calibration.

10.5. Nature's Heuristics Do Not Scale

The biological argument is compelling for biological scales but may not transfer to the complexity of human cognitive tasks, natural language, or general-purpose reasoning. FIT does not claim that heuristic navigation is sufficient for all tasks—it claims that heuristic navigation implements a special case of the gradient law, and that the gradient law provides a structural foundation for all navigating systems. Whether FIT produces competitive performance on NLP or vision benchmarks is an open empirical question—not a theoretical defeat. The theoretical contribution is the structural characterization; the empirical performance question requires experimental validation.

11. Discussion: Limits, Roadmap, and Open Questions

11.1. What FIT Does Not Claim

FIT does not claim that data is useless—data remains essential for estimating P-D fields from observations. FIT does not claim that statistical learning should be abandoned—statistical methods are the primary tools for P-D estimation. FIT does not claim that current deep learning systems do not work—they achieve remarkable performance on many tasks. FIT claims only that their performance is bounded by their training objective (behavioral correlation), that this bound is not fundamental, and that reorienting toward structural field estimation can improve efficiency, generalization, and interpretability.

11.2. The Experimental Roadmap

Empirical validation of FIT requires the following sequence:

Phase 1—Proportionality validation: For each of five domains (graph navigation, ACO on TSP, gradient descent on synthetic landscapes, electrical network routing, sensor networks), measure P and D independently, compute $F = P/D$, and test $R^2 \geq 0.80$ against measured navigability. Pre-register thresholds before observing data. Report all results including failures.

Phase 2—P-D field estimation: Train separate P-net and D-net architectures on the measurement protocols described in Section 6.3. Compare P-D estimation accuracy against naive single-field baselines.

Phase 3—FIT vs. behavioral learning: On the same five domains, train FIT models and behavioral-pattern models with matched compute budgets. Compare: (a) performance as a function of training set size (data efficiency prediction), (b) generalization after structural interventions (causal generalization prediction), (c) explanation fidelity (interpretability prediction).

Phase 4—E(t) validation: Track the coefficient of variation of the Freedom field $E(t) = \text{Var}[F]/\text{Mean}[F]$ in real-time during ACO, PSO, and gradient descent runs. Verify that convergence events correlate with $E(t)$ crossing 1.0. Report false positive rates. All code, data, and results will be published open-source. Pre-registration of hypotheses and thresholds is mandatory before any experiment is run.

11.3. The Deeper Question

Behind the technical proposals lies a deeper question: Is the universe fundamentally navigational? Does $F = P/D$ govern not just motion and information flow, but the structure of possibility itself? The AFI framework's first thesis—that Freedom is the irreducible first condition for any coherent system—is a claim about ontology, not just engineering [19,23]. We do not claim to have resolved this question. The technical proposals in this paper—the gradient law, FIT, the $E(t)$ metric—are testable regardless of whether the ontological thesis is correct. The value of proposing the ontological thesis is not that it is obviously true, but that it is productive: it generates testable hypotheses, it unifies apparently disparate phenomena, and it reframes the engineering problem from behavioral correlation to structural field estimation. This is what theories are for.

12. Conclusions

The machine learning community has debated, productively, whether to optimize models or data. We propose that this debate, while important, operates within a shared assumption that deserves scrutiny: that intelligence is the extraction of statistical patterns from behavioral records. We argue that this conflates measurement (data) with mechanism (structural fields), and that a deeper goal—learning the Perception and Distortion fields that generate behavior—may be both more efficient and more powerful.

The Law of Freedom, $F = P/D$, provides the structural foundation for this alternative. Its uniqueness under scale-covariance, monotonicity, and dimensional separability makes it falsifiable in a precise way. Its gradient law $dx/dt = -P(x) \cdot \nabla D(x)$ is structurally identical to Fick's diffusion, chemotaxis, Ohm's law, and gradient descent—revealing a structural unity that purely statistical models cannot explain. The observation that natural systems (ant colonies, immune cells, foraging animals) learn without labeled datasets, by calibrating Perception and Distortion fields through stigmergy, receptor adaptation, and heuristic compression, provides a biological existence proof for an alternative learning paradigm.

Freedom Intelligence Training (FIT) operationalizes this alternative: train P-nets and D-nets on field measurements (discriminability and traversal cost), then predict navigation by applying the gradient law. FIT predicts exponentially better data efficiency, improved causal generalization, natural interpretability, and a quantifiable exploration–exploitation transition metric. All five predictions are falsifiable with pre-registered thresholds.

We do not claim that the data paradigm has failed. Its achievements—language models, protein folding predictors, autonomous perception systems—are genuine and remarkable. We claim only that its successes are bounded by its training objective, and that the bound is not fundamental. What is fundamental, we propose, is structure: the Perception and Distortion fields that govern where every system in the universe can go, has gone, and will go next.

If not freedom, then what?

Author Contributions: Conceptualization, G.M.; methodology, G.M.; formal analysis, G.M.; investigation, G.M.; writing—original draft preparation, G.M.; writing—review and editing, G.M. The author has read and agreed to the published version of the manuscript.

Funding: This research was supported by the Portuguese Foundation for Science and Technology (FCT) through Project 2025.00020.AIVLAB.DEUCALION, providing access to the Deucalion supercomputer at the National Advanced Computing Centre (MACC), Guimarães, Portugal.

Institutional Review Board Statement: Not applicable.

Informed Consent Statement: Not applicable.

Data Availability Statement: This article develops a theoretical and conceptual framework. No empirical datasets are used. The experimental roadmap in Section 11.2 specifies the data collection protocols for future validation work, which will be published open-source with full reproducibility materials upon completion.

Acknowledgments: During the preparation of this work, the author used Claude (Anthropic) for literature search assistance, mathematical verification, and manuscript preparation. The author reviewed and edited all content and takes full responsibility for the content of the publication.

Conflicts of Interest: The author declares no conflicts of interest. The funders had no role in the design of the study; in the collection, analyses, or interpretation of data; in the writing of the manuscript; or in the decision to publish the results.

References

1. Jakubik, J.; Vossing, M.; Kühl, N.; Walk, J.; Satzger, G. Data-centric artificial intelligence. *Business & Information Systems Engineering* 2024, 66, 249–263.
2. Ng, A. MLOps: From model-centric to data-centric AI. In *Proceedings of the DeepLearning.AI Symposium*, 2022.
3. Mirzadeh, S.I.; Alizadeh, M.; Mehrabian, O.; Del Ser, J. GSM-Symbolic: Understanding the limitations of mathematical reasoning in large language models. *arXiv* 2025, arXiv:2410.05229.
4. Kejriwal, M.; Ratner, A.; Fabbri, A. Challenges, evaluation and opportunities for open-world learning. *Nature Machine Intelligence* 2024, 6, 580–588.
5. Zecevic, M.; Willig, M.; Dhami, D.S.; Kersting, K. Causal parrots: Large language models may talk causality but are not causal. *Transactions on Machine Learning Research* 2023, TMLR-2023.
6. Melo, G. The Architecture of Freedom Intelligence: A Design Theory of Path Availability. *Independent Research Manuscript* 2026. (Preprint, SSRN 6304936).
7. Berg, H.C.; Brown, D.A. Chemotaxis in *Escherichia coli* analysed by three-dimensional tracking. *Nature* 1972, 239, 500–504.
8. Fick, A. Über Diffusion. *Annalen der Physik* 1855, 170, 59–86.
9. LeCun, Y. A path towards autonomous machine intelligence. *OpenReview* 2022. <https://openreview.net/pdf?id=BZ5a1r-kVsf>.
10. Shannon, C.E. A mathematical theory of communication. *Bell System Technical Journal* 1948, 27, 379–423.
11. Raissi, M.; Perdikaris, P.; Karniadakis, G.E. Physics-informed machine learning. *Nature Reviews Physics* 2021, 3, 422–440.
12. Dorigo, M.; Gambardella, L.M. Ant colony system: A cooperative learning approach to the traveling salesman problem. *IEEE Trans. Evol. Comput.* 1997, 1, 53–66.
13. Gigerenzer, G.; Todd, P.M. *Simple Heuristics That Make Us Smart*. Oxford University Press: Oxford, UK, 1999.
14. Richens, J.G.; Lee, C.M.; Johri, S. Robust agents learn causal world models. *Proceedings of ICLR* 2024.
15. Goyal, A.; Bengio, Y. Inductive biases for deep learning of higher-level cognition. *Proceedings of the Royal Society A* 2022, 478, 20210068.
16. Hafner, D.; Lillicrap, T.; Ba, J.; Norouzi, M. DreamerV3: Mastering diverse domains through world models. *Nature* 2025.
17. Del Ser, J.; Maes, F.; Goecks, J. World models in artificial intelligence: Sensing, learning, and reasoning like a child. *arXiv* 2025, arXiv:2503.15168.
18. Pearl, J. *Causality: Models, Reasoning, and Inference*, 2nd ed.; Cambridge University Press: Cambridge, UK, 2009.

19. Gibson, J.J. *The Ecological Approach to Visual Perception*; Houghton Mifflin: Boston, MA, USA, 1979.
20. Sunstein, C. *On Freedom*; Princeton University Press: Princeton, NJ, USA, 2019.
21. Melo, G. A Structural Law of Swarm Intelligence: Path-Theoretic Unification of Navigability, Motion, and Direction. Preprint, SSRN 2026.
22. Schölkopf, B.; Locatello, F.; Bauer, S.; Ke, N.R.; Kalchbrenner, N.; Goyal, A.; Bengio, Y. Toward causal representation learning. *Proceedings of the IEEE* 2021, 109, 612–634.
23. Ashby, W.R. *An Introduction to Cybernetics*; Chapman & Hall: London, UK, 1956.
24. Friston, K. The free-energy principle: A unified brain theory? *Nature Reviews Neuroscience* 2010, 11, 127–138.
25. Hillier, B.; Hanson, J. *The Social Logic of Space*; Cambridge University Press: Cambridge, UK, 1984.
26. Kennedy, J.; Eberhart, R.C. Particle swarm optimization. In *Proceedings of the ICNN 1995, Perth, Australia, 27 November–1 December 1995; Volume 4*, pp. 1942–1948.
27. Sutton, R.S.; Barto, A.G. *Reinforcement Learning: An Introduction*, 2nd ed.; MIT Press: Cambridge, MA, USA, 2018.
28. Zhu, X.; Ghahramani, Z. Data-centric AI: A survey. *Journal of Intelligent Information Systems* 2024.
29. Lobo, J.L.; Del Ser, J. Can transformative AI shape a new age for our civilization? *arXiv* 2024, arXiv:2412.08273.
30. Maturana, H.R.; Varela, F.J. *Autopoiesis and Cognition: The Realization of the Living*; D. Reidel: Dordrecht, The Netherlands, 1980.
31. Meadows, D.H. *Thinking in Systems: A Primer*; Chelsea Green: White River Junction, VT, USA, 2008.
32. Sen, A. *Development as Freedom*; Knopf: New York, NY, USA, 1999.
33. Goyal, A.; Lamb, A.; Hoffmann, J. Recurrent independent mechanisms. In *Proceedings of ICLR* 2021.
34. Bengio, Y.; Courville, A.; Vincent, P. Representation learning: A review and new perspectives. *IEEE TPAMI* 2021, 35, 1798–1828.
35. Baesens, B.; Verbeke, W. On the design and evaluation of machine learning models for prediction tasks in engineering. *Expert Systems with Applications* 2021.

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.