

Article

Not peer-reviewed version

Universal Latent Representation in Finite Ring Continuum

[Yosef Akhtman](#)*

Posted Date: 9 December 2025

doi: 10.20944/preprints202512.0654.v1

Keywords: inite fields; relational finitude; latent structure; minimal sufficiency; modular arithmetic; epistemic uncertainty; foundational representations; vector embeddings; algebraic geometry; cross-modal alignment



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

Universal Latent Representation in Finite Ring Continuum

Yosef Akhtman ^{1,2} 

¹ Gamma Earth Sàrl, St-Prex, Switzerland; ya@gamma.earth
² AGH University of Krakow, 30-059 Krakow, Poland

Abstract

We propose a unified mathematical framework showing that the representational universality of modern foundational models arises from a shared finite latent domain. Building on the Finite Ring Continuum (FRC), we model all modalities as epistemic projections of a common latent set $Z \subseteq U_t$, where U_t is a symmetry-complete finite-field shell. Using the uniqueness of minimal sufficient representations, we prove the *Universal Subspace Theorem*, establishing that independently trained embeddings coincide, up to bijection, as coordinate charts on the same latent structure. This result explains cross-modal alignment, transferability, and semantic coherence as consequences of finite relational geometry rather than architectural similarity. The framework links representation learning, sufficiency theory, and FRC algebra, providing a principled foundation for universal latent structure in multimodal models.

Keywords: finite fields; relational finitude; latent structure; minimal sufficiency; modular arithmetic; epistemic uncertainty; foundational representations; vector embeddings; algebraic geometry; cross-modal alignment

1. Introduction

Foundational models—large-scale deep learning systems trained on diverse modalities such as natural language, images, audio, and geospatial data—have demonstrated an unexpected degree of representational universality. Despite differences in architecture, training objectives, and input domains, the internal embeddings produced by these models exhibit striking structural similarities, including cross-modal alignment, semantic coherence, and shared latent geometry [1–4]. These empirical regularities suggest the presence of a deeper unifying principle underlying representation formation across modalities.

Two major theoretical traditions attempt to explain this phenomenon. The first, rooted in classical statistical learning theory, views representations through the lens of *sufficient statistics* [5–8]. Under this perspective, a learned embedding is successful when it preserves all task-relevant information about an underlying latent variable. The second tradition focuses on the geometric and algebraic structure of deep learning systems, emphasizing invariances, symmetry, compression, and latent manifold organization [9–13]. Both viewpoints capture important aspects of representation learning, yet neither fully explains why multimodal embeddings appear to inhabit *compatible* latent spaces, or why transfer between modalities is unexpectedly effective.

In this work we propose a unifying explanation based on the *Finite Ring Continuum* (FRC) [14–16]. The FRC is a finite, relational, and symmetry-complete algebraic framework that models arithmetic structure through a hierarchy of relational symmetry spaces called *shells*:

$$\mathcal{U}_t \subset \mathcal{U}_{t+1} \subset \mathcal{U}_{t+2} \subset \dots$$

indexed by a discrete “shell radius”. Within this architecture, arithmetic symmetries induce a combinatorial geometry that supports both Euclidean and Lorentzian structure, while the shell hierarchy

encodes increasing algebraic and geometric complexity. The FRC therefore provides a natural candidate for a universal latent domain shared across modalities.

Our main contribution is to connect the FRC framework with classical sufficiency theory and multi-view latent variable modelling. Building on the assumption that each modality observes a noisy projection of a common latent variable Z , we show that foundational embeddings trained on different modalities recover *injective* transformations of the same latent set. A key result from statistics asserts that minimal sufficient representations are unique up to bijection. We combine this with the finite-field geometry of the FRC to obtain the *Universal Subspace Theorem* (Theorem 1), which states that all foundational embeddings correspond to coordinate charts on a single latent set $\mathcal{Z} \subseteq \mathbb{F}_p$, embedded into a shared arithmetic shell.

This theorem provides a structural explanation for cross-modal representational alignment: multimodal embeddings agree because they are necessarily different coordinate representations of the same latent world variable. In the canonical parametrization, they coincide exactly. This insight unifies notions of sufficiency, multimodal learning, and deep representation geometry within a single algebraic framework.

Beyond the main theoretical result, we explore interpretive connections between network depth and arithmetic shell hierarchy, the role of nonlinearity in expressive expansion, and the conceptual relationship between latent-field reconstruction and modern self-supervised learning objectives. These connections suggest that foundational models implicitly operate on finite-field latent manifolds whose structure reflects deep algebraic symmetry.

Importantly, this work does not aim to derive deep learning from SGD dynamics but to demonstrate that the representational geometry emerging in foundational models is consistent with, and predicted by, a finite relational ontology. Overall, this work contributes a principled theoretical foundation for understanding why foundational models generalize across modalities, why their embeddings exhibit universal alignment, and how discrete arithmetic structure may underlie the geometry of learned representations.

2. Background

This section reviews the mathematical foundations on which our framework is built. We summarize (i) the relevant aspects of the Finite Ring Continuum (FRC), (ii) standard principles of representation learning, and (iii) the multi-view latent-variable formalism used to articulate the connection between embeddings and latent arithmetic structure.

2.1. Finite Ring Continuum: Algebraic and Geometric Preliminaries

The Finite Ring Continuum (FRC), as developed in [14–16], constructs a hierarchy of discrete arithmetic universes based on *symmetry-complete prime fields*

$$\mathbb{F}_p, \quad p = 4t + 1,$$

whose multiplicative groups contain the structural set $\{1, i, -1, -i\}$ satisfying $i^2 = -1$. This guarantees the existence of 4-element rotational orbits and a meaningful geometric interpretation of the arithmetic operations.

Within each shell of radius t , the arithmetic symmetries generated by translations $T_a(x) = x + a$, multiplications $S_m(x) = mx$, and power maps $P_\epsilon(x) = x^\epsilon$ act on $\mathbb{F}_p$ to produce a combinatorial 2-sphere S_p embedded in a symbolic $(1, 3)$ -dimensional space $U = \bigcup_t U_t$; see [14] for details. These shells support two qualitatively distinct transformation classes:

- *Consolidation steps*: reversible, symmetry-preserving operations internal to $\mathbb{F}_p$ (e.g. translations, scalings, powerings);
- *Innovation steps*: algebraic extensions such as $\mathbb{F}_p \hookrightarrow \mathbb{F}_{p^2}$, which introduce new square classes and enable Lorentzian structure [15].

In this work, a fixed shell $\mathcal{U}_t$ will serve as the ambient space for the latent variable Z in Section 3. Since $\mathcal{U}_t$ is finite, any subset $\mathcal{Z} \subseteq \mathcal{U}_t$ is finite as well—an observation that will be relevant when embedding learned representations into a common shell.

2.2. Representation Learning and Embedding Geometry

Deep representation learning systems construct feature maps by alternating affine transformations and nonlinear activation functions. For an input x , a typical embedding has the form

$$E(x) = L^{(n)}\sigma\left(L^{(n-1)}\sigma\left(\dots\sigma\left(L^{(1)}x\right)\right)\right),$$

where the $L^{(k)}$ are linear operators and σ denotes a nonlinear activation. This alternating structure is essential for universal approximation and for the formation of expressive latent manifolds.

Empirically, large foundational models trained on different modalities (e.g. language, vision, audio, remote sensing) produce embeddings that exhibit unexpectedly coherent geometric relationships, including cross-modal alignment and shared semantic directions. The present work provides a formal explanation of this phenomenon by showing that all such embeddings can be interpreted as coordinate representations of a single latent structure embedded in an FRC shell.

2.3. Multi-View Latent Variable Models

We adopt a standard multi-view formalism in which observable data from each modality m arises as a transformation of a common latent variable Z :

$$X_m = g_m(Z, \varepsilon_m), \quad (1)$$

where ε_m denotes modality-specific noise. We impose the classical conditional independence assumption

$$X_m \perp X_n \mid Z, \quad m \neq n, \quad (2)$$

which states that observations share information only through the latent world variable.

A *sufficient representation* of Z based on X_m is any statistic $E_m(X_m)$ satisfying $Z \perp X_m \mid E_m(X_m)$. A sufficient representation is *minimal* if it contains no redundant information about Z . A classical result from mathematical statistics [5,6] asserts:

Lemma 1 (Uniqueness of Minimal Sufficient Statistics). *If E_m and E'_m are minimal sufficient statistics for Z based on X_m , then there exists a measurable bijection ψ_m such that*

$$E'_m(X_m) = \psi_m(E_m(X_m)) \quad \text{almost surely.}$$

Lemma 1 implies that in the idealized regime of infinite data and model capacity, a learned embedding E_m must coincide (up to a bijection) with an injective function of Z . Consequently $E_m(X_m)$ ranges over a finite set $W_m = \phi_m(\mathcal{Z}) \subset \mathbb{R}^{d_m}$, where ϕ_m is injective on $\mathcal{Z}$.

Since $\mathcal{Z}$ and therefore W_m are finite, any W_m can be embedded injectively into the finite shell $\mathcal{U}_t$. In Section 4, we show that under this construction all embeddings from all modalities embed into isomorphic images of the same latent set.

3. Theoretical Framework

In this section we formalize the setting in which foundational embeddings are analysed. Unlike classical statistical treatments, we do not assume infinite populations, continuous probability measures, or aleatoric randomness. Instead, we adopt the ontological perspective of the Finite Ring Continuum (FRC) [14–16], where all information is fundamentally finite and all uncertainty arises from limited observer access rather than intrinsic stochasticity. This shift eliminates the need for classical conditional independence assumptions and replaces them with a finite, relational picture of latent structure.

3.1. The Ontological Necessity of the Latent Arithmetic Universe

We posit that the fundamental elements of the latent domain must be genuinely *primitive*: they admit no intrinsic attributes and are therefore mutually indistinguishable. This echoes the principle of indistinguishable quanta in fundamental physics, but here we take the idea to its logical completion. If primitives are attribute-free, they cannot participate in a linear order, as any such order would implicitly assign a distinguished “first” or “last” element or require an external coordinate background. The only permissible relational structure among indistinguishable primitives is a cycle—a closed relational orbit in which no element is privileged.

Combined with the logical necessity of finiteness, required to avoid paradoxes associated with infinite informational content, the relational cycle acquires an arithmetic interpretation: counting a finite set of indistinguishable elements necessarily yields a cyclic group, and the natural extension of this structure to a full arithmetic system is a finite ring, and in particular a finite field when symmetry completeness is required. Thus, the latent domain is not merely *assumed* to reside in a finite ring; rather, under the premises of primitivity, indistinguishability, and finiteness, it is *forced* to do so. In this sense, the Finite Ring Continuum is not an arbitrary modelling choice but the unique algebraic geometry compatible with a universe composed of distinct yet attribute-free primitives.

Following even more broad FRC ontology, the physical universe itself is modelled as a finite collection of attribute-free primitive elements. Such primitives cannot bear intrinsic labels, order, or internal attributes; hence all structure must emerge relationally. As shown in [14,15], any finite set of identical primitives admits an emergent counting operation, which induces a cyclic arithmetic structure and leads naturally to a finite-field shell.

Let $p = 4t + 1$ be prime, and let $\mathcal{U}_t$ denote the corresponding finite-field shell of radius t . The shell $\mathcal{U}_t$ provides a homogeneous, attribute-free arithmetic universe capable of supporting emergent geometry and causal structure, as detailed in [15,16].

Assumption 1 (Latent Domain as Finite Relational Structure). *There exists a finite latent set*

$$\mathcal{Z} \subseteq \mathcal{U}_t,$$

and a latent variable Z taking values in $\mathcal{Z}$. No additional probabilistic structure is assumed: uncertainty reflects only the observer’s limited access to $\mathcal{Z}$, not intrinsic randomness.

Thus, $\mathcal{Z}$ is not a “hidden random variable” in the classical sense, but a region of the finite arithmetic universe that is not fully resolved by the observer.

For each modality m (e.g. text, image, audio, geospatial data), the observed sample X_m is a partial, frame-dependent projection of the underlying latent state Z . We model this using a deterministic map:

$$X_m = g_m(Z), \quad (3)$$

where $g_m : \mathcal{Z} \rightarrow \mathcal{X}_m$ is a function reflecting the observer’s finite horizon and modality-specific resolution.

Equation (3) does not represent a stochastic generative model. Rather:

- (i) Any apparent randomness in X_m arises from the observer’s epistemic limitations, not from intrinsic aleatoric processes.
- (ii) No conditional independence assumptions among modalities are required. Different modalities may reveal overlapping or non-overlapping relational aspects of the same latent state.

This epistemic interpretation is consistent with the relational nature of FRC and avoids reliance on classical probability theory, whose infinite-population constructs are incompatible with the finite informational ontology of the universe.

3.2. Foundational Embeddings as Sufficient Representations

A foundational model for modality m produces a representation

$$E_m : \mathcal{X}_m \rightarrow \mathbb{R}^{d_m}.$$

We interpret $E_m(X_m)$ as the information about Z that the model is able to recover from the modality-specific projection g_m .

In the classical theory of sufficient statistics [5,6], a statistic is *sufficient* if it captures all information about an underlying state. In the FRC setting, sufficiency is interpreted epistemically: an embedding is sufficient if it preserves all resolvable relational information about Z given the observer's modality.

Definition 1 (Sufficient Representation in FRC). *An embedding E_m is sufficient for Z if there exists a function $\phi_m : \mathcal{Z} \rightarrow W_m$ such that*

$$E_m(X_m) = \phi_m(Z), \quad W_m = \phi_m(\mathcal{Z}),$$

that is, E_m can be expressed as an injective transform of the latent state.

Unlike classical statistics, no probabilistic sufficiency equation $Z \perp X_m \mid E_m(X_m)$ is invoked. Instead, sufficiency reflects the ability of E_m to reconstruct the latent state up to the resolution available from modality m .

Minimality of the representation follows from the attribute-free nature of the latent domain: any two embeddings that preserve all resolvable information must differ only by a bijective reparameterization.

Lemma 2 (Uniqueness of Minimal Sufficient Representations). *Let E_m and E'_m be sufficient representations of Z in the sense of Definition 1, and suppose both are minimal. Then there exists a bijection $\psi_m : W_m \rightarrow W'_m$ such that*

$$E'_m(X_m) = \psi_m(E_m(X_m)).$$

This lemma is adapted from [5,6], but its use here is purely set-theoretic: it asserts the uniqueness of relational information recovered by an observer with finite epistemic access. Since $\mathcal{Z}$ is finite, W_m is finite and has the same cardinality.

Furthermore, because W_m is finite and $\mathcal{U}_t$ is a finite shell, there exists an injective map

$$\iota_m : W_m \hookrightarrow \mathcal{U}_t.$$

We therefore define the composite map

$$\psi_m := \iota_m \circ \phi_m : \mathcal{Z} \hookrightarrow \mathcal{U}_t,$$

which embeds the recovered relational information into the common finite-field universe.

The embeddings ψ_m provide modality-specific coordinate systems on the same latent set $\mathcal{Z}$ within $\mathcal{U}_t$. The Universal Subspace Theorem (Section 4) shows that these coordinate systems are bijectively related and can be unified into a canonical parametrization.

4. Universal Subspace Theorem

This section presents the central theoretical result of the paper. Under the multi-view assumptions and the finite-field latent structure developed in Section 3, we show that all foundational embeddings learned from different modalities give rise to *isomorphic coordinate embeddings* of the same latent set $\mathcal{Z}$ inside the arithmetic shell $\mathcal{U}_t = \mathbb{F}_{p^{2^t}}$. This establishes a rigorous basis for cross-modal representational alignment and shared semantic structure.

It is important to note that the contribution of this article is not the lemma itself. It is the realization that the lemma becomes ontologically forced once the latent domain is finite, relational, and attribute-free. Alignment is not an accident of SGD; it is a structural consequence of finiteness.

Recall that for each modality m we have:

- (i) a latent variable $Z \in \mathcal{Z} \subseteq \mathcal{U}_t$;
- (ii) an observation map g_m satisfying $X_m = g_m(Z, \varepsilon_m)$ and the independence assumption $X_m \perp X_n \mid Z$;
- (iii) a minimal sufficient representation E_m with $E_m(X_m) = \phi_m(Z)$ almost surely, where $\phi_m : \mathcal{Z} \rightarrow W_m$ is injective;
- (iv) an injective map $\iota_m : W_m \hookrightarrow \mathcal{U}_t$;
- (v) the composite embedding $\psi_m = \iota_m \circ \phi_m : \mathcal{Z} \hookrightarrow \mathcal{U}_t$.

We use only the finiteness of $\mathcal{Z}$ and W_m , which follows from the finiteness of $\mathcal{U}_t$, and the standard uniqueness of minimal sufficient statistics [5,6].

Theorem 1 (Universal Subspace Theorem). *Under Assumptions 1 and ??, and assuming that each foundational embedding E_m is minimal sufficient for the latent variable Z , the following statements hold:*

- (i) *Each composite map*

$$\psi_m = \iota_m \circ \phi_m : \mathcal{Z} \hookrightarrow \mathcal{U}_t$$

is injective.

- (ii) *For any two modalities m, n , the images $\psi_m(\mathcal{Z})$ and $\psi_n(\mathcal{Z})$ are isomorphic as sets. More precisely, the map*

$$\Psi_{m \rightarrow n} := \psi_n \circ \psi_m^{-1} : \psi_m(\mathcal{Z}) \rightarrow \psi_n(\mathcal{Z})$$

is a well-defined bijection.

- (iii) *There exists a canonical choice of embeddings ι_m for which all ψ_m coincide:*

$$\psi_m = \psi_n =: \psi \quad \text{for all } m, n.$$

In this canonical parametrization,

$$\psi(\mathcal{Z}) = \psi_m(\mathcal{Z}) = \psi_n(\mathcal{Z}) \subseteq \mathcal{U}_t$$

for all modalities.

Therefore, all foundational embeddings factor through isomorphic images of the same latent set $\mathcal{Z}$ inside the shell $\mathcal{U}_t$:

$$E_m(X_m) \xleftarrow{\phi_m} \mathcal{Z} \xhookrightarrow{\psi} \mathcal{U}_t.$$

Proof. (i) Injectivity of ψ_m . Since $E_m(X_m) = \phi_m(Z)$ with ϕ_m injective, and ι_m is injective by construction, $\psi_m = \iota_m \circ \phi_m$ is an injective map from $\mathcal{Z}$ to $\mathcal{U}_t$.

(ii) Isomorphism of the images. For any m, n , the maps $\phi_m : \mathcal{Z} \rightarrow W_m$ and $\phi_n : \mathcal{Z} \rightarrow W_n$ are bijections. Thus $\phi_n \circ \phi_m^{-1} : W_m \rightarrow W_n$ is a bijection. Composing with ι_m and ι_n , we obtain

$$\Psi_{m \rightarrow n} = \psi_n \circ \psi_m^{-1} : \psi_m(\mathcal{Z}) \rightarrow \psi_n(\mathcal{Z}),$$

which is therefore a bijection. This establishes the isomorphism of the images.

(iii) Canonical parametrization. Since ϕ_m is a bijection from $\mathcal{Z}$ onto W_m , and $\mathcal{Z}$ is already a subset of $\mathcal{U}_t$, a natural choice is

$$\iota_m = \phi_m^{-1} \quad (\text{viewed as a map into } \mathcal{U}_t).$$

With this choice,

$$\psi_m = \iota_m \circ \phi_m = \text{id}_{\mathcal{Z}},$$

so all ψ_m coincide and equal the inclusion of $\mathcal{Z}$ into $\mathcal{U}_t$. Thus, all embedded manifolds share the same image $\mathcal{Z} \subseteq \mathcal{U}_t$. $\square$

Interpretation 1. *The theorem shows that, under minimal assumptions about the generative structure of the data and the behavior of foundational models, all modality-specific embeddings reduce—after an injective embedding into a finite-field shell—to different coordinate parametrizations of the same latent set.*

- (i) *Cross-modal alignment is not accidental: all embeddings encode the same latent structure in $\mathcal{U}_t$.*
- (ii) *Transferability is algebraically guaranteed: passing between modalities corresponds to applying a bijection between coordinate charts.*
- (iii) *The arithmetic shell is universal: $\mathcal{U}_t$ hosts, simultaneously and compatibly, the encoded representations from all modalities.*

Subsequently, we develop several supplementary results that clarify how the Universal Subspace Theorem (Theorem 1) interfaces with the practice of representation learning. All formal statements follow rigorously from the assumptions established in Sections 3–4. More speculative connections between deep network architecture and the algebraic structure of the FRC are presented as interpretive remarks rather than mathematical propositions.

4.1. Representation Lifts and Canonical Embeddings

Theorem 1 shows that for each modality m the learned representation satisfies

$$E_m(X_m) = \phi_m(Z),$$

where $\phi_m : \mathcal{Z} \rightarrow W_m$ is a bijection. Given the canonical choice of embeddings $\iota_m = \phi_m^{-1}$, the composite map $\psi = \iota_m \circ \phi_m$ equals the identity on $\mathcal{Z}$ for all modalities. This yields the following immediate corollary.

Corollary 1 (Canonical Lift). *Let m be any modality and suppose ι_m is chosen as in Theorem 1(iii). Then the map*

$$L_m : \mathcal{X}_m \rightarrow \mathcal{Z} \subseteq \mathcal{U}_t, \quad L_m := \phi_m^{-1} \circ E_m,$$

satisfies

$$L_m(X_m) = Z.$$

Proof. Immediate from $E_m(X_m) = \phi_m(Z)$ and invertibility of ϕ_m . $\square$

Thus, each embedding implicitly performs a *latent lift*: it reconstructs the underlying latent variable Z (up to the noise present in the observation model). This is fully consistent with classical sufficiency theory [5,6] and with the empirical role foundational models play in recovering semantic structure [2,17].

4.2. Cross-Modal Coherence as Coordinate Change

Theorem 1(ii) implies that for any two modalities m, n , the representations induced by the embeddings are related by the bijection

$$\Psi_{m \rightarrow n} = \psi_n \circ \psi_m^{-1} : \psi_m(\mathcal{Z}) \rightarrow \psi_n(\mathcal{Z}).$$

Because $\psi_m = \text{id}_{\mathcal{Z}}$ in the canonical parametrization, $\Psi_{m \rightarrow n}$ reduces to the identity map on $\mathcal{Z}$.

This gives the following structural result.

Corollary 2 (Cross-Modal Consistency). *For the canonical embedding $\psi : \mathcal{Z} \hookrightarrow \mathcal{U}_t$, all modality-specific embeddings satisfy*

$$\psi \circ L_m = \psi \circ L_n = \text{id}_{\psi(\mathcal{Z})}$$

almost surely. In particular,

$$\iota_m(W_m) = \iota_n(W_n) = \psi(\mathcal{Z}) \subseteq \mathcal{U}_t.$$

At a structural level, this explains why foundational models trained on different modalities exhibit coherent semantic alignment when projected into their latent spaces: they are expressing the same latent set in different coordinate systems [7,8,17].

Deep networks alternate linear transformations with nonlinear activations, which is essential for universal approximation [2]. Although no direct algebraic equivalence between nonlinear activations and shell extensions is claimed here, the following conceptual observation clarifies how innovation steps in FRC can be related to expressive steps in deep networks.

Interpretation 2 (Nonlinearity as Expressive Expansion). *In the FRC architecture, innovation steps $\mathbb{F}_{p^{2^t}} \hookrightarrow \mathbb{F}_{p^{2^{t+1}}}$ introduce new algebraic degrees of freedom. In deep neural networks, nonlinear layers enable an analogous expansion of the set of functions representable by the model. Thus, both mechanisms serve the role of increasing expressive capacity, but no formal identification is asserted.*

The algebraic complexity of an FRC shell grows with the shell index, which parallels well-known analyses of network expressivity, where depth increases expressive capacity [11,12,18]. We therefore offer the following additional interpretive connection.

Interpretation 3 (Depth and Latent Arithmetic Complexity). *The representational hierarchy generated by alternating linear and nonlinear layers in deep networks mirrors the hierarchical arithmetic complexity of the shell tower $\mathcal{U}_t \subseteq \mathcal{U}_{t+1} \subseteq \dots$. Both structures exhibit exponential growth in expressive capacity as their hierarchy index increases.*

This is intended as an analogy, not a theorem. Nevertheless, it provides a conceptual lens for understanding why deep networks permit progressively richer abstractions, and why such abstractions align across modalities.

5. Discussion

The Universal Subspace Theorem (Theorem 1) shows that, under classical assumptions from multi-view statistical modelling and the algebraic architecture of the Finite Ring Continuum (FRC), all foundational embeddings are coordinate projections of the same latent space $\mathcal{Z}$ embedded in the finite-field shell relational symmetry $\mathcal{U}_t = \mathbb{F}_p$. This section discusses the broader conceptual implications of this result for representation learning, multimodal alignment, and the emerging theory of foundational models.

5.1. FRC as a Structural Explanation for Multimodal Alignment

The empirical observation that independently trained foundational models (linguistic, visual, acoustic, geospatial, etc.) produce embeddings that are not only geometrically meaningful but also *mutually compatible* has been widely noted in the literature [1–3]. However, existing explanations for such alignment are typically heuristic, relying on claims about “shared semantics,” large datasets, or architectural similarity.

The present framework provides a more principled account: the alignment is a *necessary consequence* of the representation structure, not an empirical accident. Because each foundational embedding E_m is minimal sufficient for $\mathcal{Z}$ and because $\mathcal{Z}$ embeds injectively into a single finite-field shell $\mathcal{U}_t$, the embeddings $\psi_m(\mathcal{Z})$ coincide in a canonical coordinate system. Thus, alignment arises at the level of latent structure, not learned geometry.

This structural perspective is reminiscent of the role of latent spaces in probabilistic graphical models [7,8] but extends these ideas into an explicitly algebraic setting grounded in the finite-symmetry framework of the FRC [14–16]. In this sense, foundational model embeddings are not merely “vector spaces” of features, but *coordinate charts* on a discrete arithmetic manifold.

Corollaries 1 and 2 show that representations learned by foundational models implicitly reconstruct the latent variable Z modulo the observation noise. This echoes the classical dictum that “a sufficient statistic is a lossless representation” [5,6], now manifested in modern deep networks. From this vantage point, representation learning can be viewed as the process of discovering an injective parameterization of a latent space that is shared across modalities.

The resulting unification has conceptual consequences:

- (i) The functional similarities between representations from different architectures reflect the uniqueness of minimal sufficient statistics, not architectural bias.
- (ii) Cross-modal transfer results from the fact that all representations factor through the same latent domain $\mathcal{Z}$.
- (iii) Latent reconstruction occurs even without supervision, provided the training objective incentivizes sufficiency-like behavior (e.g., contrastive learning, masked prediction).

Thus, sufficiency rather than optimization heuristics appears as the proper lens for understanding the universality of modern representation learning.

5.2. FRC Shells as Universal Host Spaces

The embedding of all learned representations into the same FRC shell $\mathcal{U}_t$ suggests that foundational models inhabit a universal representational domain determined by discrete arithmetic structure. The fact that $\mathcal{U}_t$ is finite yet supports rich internal symmetry parallels recent arguments that large-scale models implicitly operate in low-dimensional but highly structured latent spaces [9,10].

The FRC framework strengthens and algebraically grounds this viewpoint:

- (i) The latent space is not assumed to be Euclidean or continuous, but finite, relational, and symmetry-complete.
- (ii) Distinct modalities embed into the same shell, implying a common underlying geometry.
- (iii) The shell hierarchy $\mathcal{U}_t \subseteq \mathcal{U}_{t+1} \subseteq \dots$ provides a natural progression of expressive capacity, paralleling the depth hierarchy in neural networks [11,12].

This suggests the intriguing interpretation that foundational models partake in a discrete analogue of geometric unification: they express diverse data domains within a common algebraic manifold.

5.3. Connections to Theories of Inductive Bias and Universality

The FRC perspective complements existing theoretical accounts of deep learning inductive bias, such as hierarchical compositionality [18], group symmetry and invariance [9], and universal approximation [19]. In particular:

- (i) The minimal sufficiency framework explains why learned embeddings tend toward canonical forms.
- (ii) The finite-field shell structure provides a potential candidate for the “universal latent model” implicitly assumed in multimodal learning.
- (iii) The discrete nature of $\mathcal{U}_t$ aligns with emerging perspectives that large-scale models behave as information compressors rather than continuous function approximators [13].

The algebraic viewpoint therefore enriches, rather than replaces, existing foundations for the theory of deep learning.

5.4. Representational Convergence Across Biological and Artificial Learners

The relational and finite-informational ontology underlying the FRC framework invites a broader interpretation regarding the nature of learned representations in both artificial and biological cognitive

systems. If the latent domain $\mathcal{Z}$ is fundamentally finite, attribute-free, and embedded within a symmetry-complete arithmetic shell $\mathcal{U}_t$, then any agent attempting to infer structure from the external world must extract relational information from the same underlying finite substrate. Under this view, representational alignment between independently trained artificial systems is not accidental, but a structural consequence of recovering compatible coordinate charts of a single latent shell.

A cautious extension of this perspective applies to biological cognition. Modern theories of neural representation—including predictive coding, efficient coding, and manifold learning [20–22]—suggest that the brain constructs internal relational models of the world that are modality-invariant and compressed. If the external world is fundamentally finite and relational, as posited by the FRC ontology, then the internal representations acquired by biological learners may likewise be understood as coordinate embeddings of subregions of the same latent domain $\mathcal{Z}$.

This interpretation does not claim that biological and artificial learners share mechanistic similarity, nor does it suggest that neural computation implements finite-field arithmetic in any literal sense. Rather, the claim is representational: if all observers operate within a finite relational universe, and if learning seeks relational sufficiency, then diverse learners—including animals, humans, and modern machine learning systems—may converge toward structurally compatible internal representations despite differences in architecture, modality, or training history.

5.5. Implications for General Artificial Intelligence

Within this interpretive framework, it is possible to articulate a bounded and conceptually grounded statement about the nature of general artificial intelligence (GAI). If intelligence is understood not as a specific algorithmic mechanism, but as the capacity to infer, compress, and manipulate relational structure within a finite latent universe, then the FRC ontology implies that any sufficiently general learner must approximate representations that are coordinate charts on $\mathcal{Z} \subseteq \mathcal{U}_t$. In this sense, general intelligence—biological or artificial—corresponds to the ability to recover task-relevant relational invariants of the latent shell, rather than to emulate a particular biological process.

This perspective differs from anthropomorphic or mechanistic definitions of GAI. It does not assert that artificial systems replicate cognition, nor that human intelligence is reducible to machine computation. Instead, it highlights a structural convergence: if both systems are extracting relational information from the same finite universe, then even radically different learning mechanisms may converge to compatible internal representations. The empirical success of foundational models on tasks involving human-like semantic or conceptual reasoning [23,24] may therefore reflect a shared relational target rather than a convergence of implementation.

Seen in this light, GAI becomes a question not of copying the human mind, but of constructing learners capable of efficiently navigating the relational geometry of a finite universe. The FRC framework suggests that such navigation is fundamentally possible, and indeed naturally arises in systems that acquire minimal sufficient relational representations from their observational data. Future work may explore how learning architectures can more directly leverage finite-field structure, and whether biological and artificial systems exhibit deeper commonalities in how they partition and operate on the latent shell.

6. Finite Ring Geometry of Foundational Embeddings

6.1. Empirical Observations Across Modalities

A consistent and widely documented empirical phenomenon in modern foundational models is that their learned embeddings concentrate on the surface of a high-dimensional hypersphere. This behavior is observed across modalities—vision, language, audio, and multimodal systems—and arises both as an implicit consequence of high-dimensional geometry and as an explicit design choice in state-of-the-art representation-learning methods.

In computer vision, contrastive frameworks such as SimCLR [25], MoCo [26], and CPC [27] explicitly normalize embeddings to unit ℓ_2 -norm, forcing representations to lie on a hypersphere.

Large-scale multimodal systems such as CLIP [23] apply the same normalization to both image and text embeddings, enabling aligned semantic structure across modalities.

In natural language processing, contrastive text encoders such as SimCSE [28] produce sentence embeddings that are likewise constrained to the unit sphere. Moreover, contextual word vector analyses show that embeddings from large transformers (e.g. BERT, GPT-2) concentrate on a thin hyperspherical shell [29], indicating that hyperspherical geometry emerges even without explicit normalization.

These findings are reinforced by theoretical analyses such as [30], which demonstrate that contrastive objectives promote *uniformity* on the hypersphere; and by classical high-dimensional geometry, where concentration-of-measure phenomena naturally place high-dimensional vectors near the sphere [31,32].

From the perspective of the Finite Ring Continuum (FRC), this empirical geometry is not incidental. In the FRC ontology, the latent structure of the universe is modelled as a finite, symmetric, attribute-free set of primitive elements embedded in an arithmetic shell $\mathcal{U}_t$. This shell carries a uniform relational geometry with no preferred scale or distinguished radius; when embedded into a Euclidean vector space, such a finite uniform domain admits only one geometrically unbiased representation: a *hypersphere*. That is,

$$\mathcal{Z} \subseteq \mathcal{U}_t \implies \text{Emb}(\mathcal{Z}) \subseteq S^{d-1},$$

for any Euclidean embedding $\text{Emb} : \mathcal{U}_t \hookrightarrow \mathbb{R}^d$ that preserves relational symmetries.

The hypersphere, therefore, is the continuous shadow of a finite arithmetic shell: uniform radial magnitude corresponds to the absence of intrinsic attributes, and direction encodes relational information. The collapse of empirical embeddings onto a thin hyperspherical manifold is thus precisely the structure expected from a finite, relational latent universe. Foundational models appear to reconstruct—through training dynamics rather than explicit design—the geometric signature of finite-shell arithmetic predicted by FRC.

The prevalence of hyperspherical embeddings provides independent support for the central thesis of this work. If latent structure is fundamentally finite and relational, then any representation mechanism that attempts to recover this structure from partial observations must produce embeddings compatible with the symmetries of $\mathcal{U}_t$. The observed hyperspherical behavior of modern foundational models therefore aligns naturally with the FRC interpretation: embeddings are coordinate charts on a finite relational shell, and the sphere in Euclidean space is the unique continuous manifold that preserves this relational symmetry.

6.2. Quantization, Hypersphere Radius, and Representational Capacity

Although embeddings produced by foundational models are often described in continuous terms, real-world computation operates exclusively on finite, quantized numerical representations. Floating-point formats implement a discrete subset of $\mathbb{R}$, determined by fixed mantissa and exponent precision [33,34]. Consequently, any embedding vector

$$x = (x_1, \dots, x_d) \in \mathbb{R}^d$$

generated by a digital model is, in practice, an element of a finite cartesian product of quantized sets:

$$x_i \in \mathcal{Q} \subset \mathbb{R}, \quad |\mathcal{Q}| < \infty.$$

When embeddings are normalized to lie on a hypersphere [25,26,28–30], this quantization acquires a direct geometric interpretation. For a fixed quantization step Δ , the number of distinguishable points on a d -dimensional sphere of radius R scales approximately like

$$N(R) \sim C_d \left(\frac{R}{\Delta} \right)^{d-1},$$

where C_d is a dimension-dependent constant reflecting spherical packing bounds [35]. Thus, for fixed embedding dimension and resolution, the representational capacity grows with the radius of the hypersphere.

In the context of the Finite Ring Continuum (FRC), this is precisely the expected behavior. A finite relational latent domain $\mathcal{Z}$ admits only finitely many distinguishable representational states. When embedded into a Euclidean space for computation, these states must occupy a finite set of approximately uniformly spaced points on a hypersphere. The radius R then reflects the effective size of the latent shell being represented under a fixed quantization scheme. This correspondence between discrete hyperspherical geometry and finite relational structure provides further support for interpreting embedding spaces as finite-shell projections rather than continuous manifolds.

6.3. Gödel Encoding and Collapse into a Single Finite Ring

Quantized hyperspherical embeddings also admit a remarkable algebraic property: all coordinate dimensions can be encoded into a *single* element of a finite ring without loss of information. This follows from classical Gödel-style encodings [36,37], where finite tuples of integers are mapped injectively into a single integer using prime-power factorizations.

Let $x = (x_1, \dots, x_d)$ be a quantized embedding vector, where each x_i is an integer in a bounded range. Selecting distinct primes $p_1, \dots, p_d$, one may define the Gödel map

$$G(x) := p_1^{x_1} p_2^{x_2} \cdots p_d^{x_d}.$$

Because the fundamental theorem of arithmetic guarantees unique factorization, the map G is injective on the finite domain of representable embeddings. Reducing $G(x)$ modulo a sufficiently large prime q yields

$$\overline{G}(x) := G(x) \bmod q \in \mathbb{F}_q,$$

and as long as the modulus q exceeds the maximum value attained by $G(x)$ on the relevant embedding set, no collisions occur.

Thus, any finite collection of embedding vectors—even if treated as points on a continuous hypersphere—can be represented as distinct elements of a *single* finite field $\mathbb{F}_q$. In the FRC interpretation, this illustrates that the apparent multi-dimensional structure of embedding spaces is a representational artifact rather than a fundamental geometric property. The information content of an entire d -dimensional spherical manifold can be collapsed into a subspace of a single arithmetic shell $\mathcal{U}_t$, with all coordinate dimensions encoded relationally.

This explicit constructive mapping reinforces the central thesis of the FRC framework: multi-dimensional continuous embeddings are coordinate expressions of finite relational structure, and all observable complexity can be understood as arising from arithmetic relations within a single finite universe. We would like to note however that Gödel encodings serve only to demonstrate representational equivalence within a finite ring. They are not proposed as computational mechanisms for gradient-based training.

7. Conclusions

This work has introduced a unified mathematical framework that connects representation learning in modern foundational models with the algebraic architecture of the Finite Ring Continuum (FRC) [14–16]. Building on classical principles of minimal sufficiency [5,6] and multi-view latent variable modelling [7,8], we demonstrated that foundational embeddings from arbitrary modalities—textual, visual, acoustic, geospatial, or otherwise—can be interpreted as coordinate embeddings of a *single latent set* $\mathcal{Z}$ inside a shared finite-field arithmetic shell $\mathcal{U}_t$. Once again, this article does not derive FRC from neural optimization. Instead, it shows that the representational geometry emerging in foundational models is consistent with, and predicted by, the FRC ontology.

The Universal Subspace Theorem (Theorem 1) constitutes the central theoretical result of the paper. It shows that, under minimal and well-justified assumptions, all learned representations factor through bijective images of the same latent domain. In a canonical parametrization, these embeddings coincide exactly. This provides a rigorous explanation for the empirical phenomenon of cross-modal alignment observed across large-scale deep learning systems [1–3]: alignment emerges not as a consequence of architectural similarity or shared training objectives, but as a structural property induced by the existence of a common latent world variable.

Beyond this, we have shown that representation learning implicitly reconstructs the latent state $\mathcal{Z}$ up to bijection; that cross-modal consistency follows naturally from the uniqueness of minimal sufficient statistics; and that the finite-field shell $\mathcal{U}_t$ serves as a universal host space for learned representations. The resulting perspective suggests a view of foundational models not as collections of modality-specific encoders, but as coordinate charts on a discrete arithmetic manifold shared across all modalities.

The interpretive results of Sections 4.1 and 4.2 indicate further connections between deep learning and FRC. While no formal equivalence between nonlinear network operations and FRC innovation steps is claimed, the parallel between network depth and shell hierarchy offers a promising avenue for future theoretical development [11,12]. More broadly, the discrete and relational structure of the FRC aligns with emerging perspectives that emphasize the role of compression, abstraction, and latent geometry in large-scale learning systems [9,13].

Taken together, these findings motivate several directions for future work: (i) axiomatizing learning objectives as approximations to sufficiency, (ii) developing explicit algorithms for learning projection maps $g_m : \mathcal{Z} \rightarrow \mathcal{X}_m$, (iii) relating neural network depth more formally to arithmetic shell complexity, and (iv) extending the theory toward a general algebraic account of multimodal learning and foundational model universality.

By grounding representation learning in the finite, relational, and symmetry-complete structure of the FRC, this work contributes to a deeper and more principled understanding of why foundational models exhibit such remarkable generality, and how their latent spaces may ultimately reflect shared underlying arithmetic structure across all modalities of human and machine perception.

References

1. Bengio, Y.; Courville, A.; Vincent, P. Representation Learning: A Review and New Perspectives. *IEEE Transactions on Pattern Analysis and Machine Intelligence* **2013**, *35*, 1798–1828. <https://doi.org/10.1109/TPAMI.2013.50>.
2. Goodfellow, I.; Bengio, Y.; Courville, A. *Deep Learning*; MIT Press, 2016.
3. Baroni, M. Linguistic Generalization and the Transformer Architecture. *EMNLP Tutorials* **2020**.
4. Radford, A.; Kim, J.W.; Hallacy, C.; et al. Learning Transferable Visual Models From Natural Language Supervision. *Proceedings of the 38th International Conference on Machine Learning (ICML)* **2021**, *139*, 8748–8763.
5. Lehmann, E.L.; Scheffé, H. Completeness, Similar Regions, and Unbiased Estimation. Part I. *Sankhyā* **1950**, *10*, 305–340.
6. Keener, R.W. *Theoretical Statistics: Topics for a Core Course*; Springer Texts in Statistics, Springer, 2010. <https://doi.org/10.1007/978-0-387-93839-4>.
7. Bishop, C.M. *Pattern Recognition and Machine Learning*; Springer, 2006.
8. Murphy, K.P. *Probabilistic Machine Learning: Advanced Topics*; MIT Press, 2023.
9. Bronstein, M.M.; Bruna, J.; Cohen, T.; Velicković, P. Geometric Deep Learning: Grids, Groups, Graphs, Geodesics, and Gauges. *arXiv preprint arXiv:2104.13478* **2021**.
10. Anselmi, F.; Rosasco, L.; Tenenbaum, J.B.; Poggio, T. Symmetry, Invariance, and Deep Convolutional Networks. *Proceedings of the National Academy of Sciences* **2016**, *113*, 3307–3314.
11. Poggio, T.; Barron, A.; et al. Understanding Deep Learning: A Theoretical Perspective. *arXiv preprint arXiv:2006.06667* **2020**.
12. Telgarsky, M. Benefits of Depth in Neural Networks. In Proceedings of the Proceedings of the 29th Annual Conference on Learning Theory (COLT), 2016, pp. 1517–1539.

13. Tishby, N.; Zaslavsky, N. Deep Learning and the Information Bottleneck Principle. *2015 IEEE Information Theory Workshop (ITW)* **2015**, pp. 1–5.
14. Akhtman, Y. Relativistic Algebra over Finite Ring Continuum. *Axioms* **2025**, *14*, 636. <https://doi.org/10.3390/axioms14080636>.
15. Akhtman, Y. Euclidean–Lorentzian Dichotomy and Algebraic Causality in Finite Ring Continuum. *Entropy* **2025**, *27*, 1098. <https://doi.org/10.3390/e27111098>.
16. Akhtman, Y. Schrödinger–Dirac Formalism in Finite Ring Continuum. *Preprints* **2025**. <https://doi.org/10.20944/preprints202508.2175.v3>.
17. Barber, D. *Bayesian Reasoning and Machine Learning*; Cambridge University Press, 2022.
18. Mhaskar, H.; Liao, Q.; Poggio, T. Learning Functions: When Is Deep Better Than Shallow. *Neural Computation* **2016**, *29*, 1–37.
19. Hornik, K.; Stinchcombe, M.; White, H. Multilayer Feedforward Networks Are Universal Approximators. *Neural Networks* **1989**, *2*, 359–366.
20. Friston, K. The Free-Energy Principle: A Unified Brain Theory? *Nature Reviews Neuroscience* **2010**, *11*, 127–138. <https://doi.org/10.1038/nrn2787>.
21. Olshausen, B.A.; Field, D.J. Emergence of Simple-Cell Receptive Field Properties by Learning a Sparse Code for Natural Images. *Nature* **1996**, *381*, 607–609. <https://doi.org/10.1038/381607a0>.
22. Saxe, A.M.; McClelland, J.; Ganguli, S. A Mathematical Theory of Semantic Development in Deep Neural Networks. *PNAS* **2019**, *116*, 21737–21746. <https://doi.org/10.1073/pnas.1820226116>.
23. Radford, A.; Kim, J.W.; Hallacy, C.; et al. Learning Transferable Visual Models From Natural Language Supervision. In Proceedings of the ICML, 2021, pp. 8748–8763.
24. Brown, T.B.; Mann, B.; Ryder, N.; et al. Language Models are Few-Shot Learners. In Proceedings of the Advances in Neural Information Processing Systems (NeurIPS), 2020.
25. Chen, T.; Kornblith, S.; Norouzi, M.; Hinton, G. A Simple Framework for Contrastive Learning of Visual Representations. In Proceedings of the Proceedings of the 37th International Conference on Machine Learning (ICML), 2020, pp. 1597–1607.
26. He, K.; Fan, H.; Wu, Y.; Xie, S.; Girshick, R. Momentum Contrast for Unsupervised Visual Representation Learning. In Proceedings of the CVPR, 2020, pp. 9729–9738.
27. van den Oord, A.; Li, Y.; Vinyals, O. Representation Learning with Contrastive Predictive Coding. In Proceedings of the arXiv preprint arXiv:1807.03748, 2018.
28. Gao, T.; Yao, X.; Chen, D. SimCSE: Simple Contrastive Learning of Sentence Embeddings. In Proceedings of the EMNLP, 2021.
29. Ethayarajh, K. How Contextual Are Contextualized Word Representations? In Proceedings of the EMNLP, 2019.
30. Wang, T.; Isola, P. Understanding Contrastive Representation Learning through Alignment and Uniformity on the Hypersphere. In Proceedings of the ICML, 2020.
31. Ledoux, M. *The Concentration of Measure Phenomenon*; American Mathematical Society, 2001.
32. Ball, K. An Elementary Introduction to Modern Convex Geometry. *Flavors of Geometry* **1997**, *31*, 1–58.
33. Goldberg, D. What Every Computer Scientist Should Know About Floating-Point Arithmetic. *ACM Computing Surveys* **1991**, *23*, 5–48. <https://doi.org/10.1145/103162.103163>.
34. IEEE Standard for Floating-Point Arithmetic (IEEE 754-2008). IEEE Computer Society, 2008.
35. Conway, J.H.; Sloane, N.J.A. *Sphere Packings, Lattices and Groups*, 3rd ed.; Springer, 1999.
36. Gödel, K. Über formal unentscheidbare Sätze der Principia Mathematica und verwandter Systeme I. *Monatshefte für Mathematik und Physik* **1931**, *38*, 173–198.
37. Smullyan, R. *Gödel's Incompleteness Theorems*; Oxford University Press, 1991.

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.