

How Language Directions Align with Token Geometry in Multilingual LLMs

JaeSeong Kim
Semyung University
Jecheon-si, Republic of Korea
mmmqp1010@gmail.com

Suan Lee
Semyung University
Jecheon-si, Republic of Korea
suanlee@semyung.ac.kr

Abstract

Multilingual LLMs demonstrate strong performance across diverse languages, yet there has been limited systematic analysis of how language information is structured within their internal representation space and how it emerges across layers. We conduct a comprehensive probing study on six multilingual LLMs, covering all 268 transformer layers, using linear and nonlinear probes together with a new Token–Language Alignment analysis to quantify the layer-wise dynamics and geometric structure of language encoding. Our results show that language information becomes sharply separated in the first transformer block ($+76.4 \pm 8.2\%$ from Layer 0→1) and remains **almost fully linearly separable** throughout model depth. We further find that the alignment between language directions and vocabulary embeddings is **strongly tied to the language composition of the training data**. Notably, Chinese-inclusive models achieve a ZH Match@Peak of 16.43%, whereas English-centric models achieve only 3.90%, revealing a $4.21 \times$ **structural imprinting** effect. These findings indicate that multilingual LLMs distinguish languages not by surface script features but by **latent representational structures shaped by the training corpus**. Our analysis provides practical insights for data composition strategies and fairness in multilingual representation learning. All code and analysis scripts are publicly available at: <https://github.com/thisiskorea/How-Language-Directions-Align-with-Token-Geometry-in-Multilingual-LLMs>.

CCS Concepts

• **Computing methodologies** → **Information extraction**; *Language resources*.

Keywords

Multilingual LLMs, Language Directions, Token Geometry, Representation Learning, Linear Probing

ACM Reference Format:

JaeSeong Kim and Suan Lee. 2026. How Language Directions Align with Token Geometry in Multilingual LLMs. In *Proceedings of The Web Conference 2026 (WWW '26)*. ACM, New York, NY, USA, 4 pages. <https://doi.org/10.1145/nnnnnnnn.nnnnnnnn>

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

WWW '26, Dubai, United Arab Emirates

© 2026 Copyright held by the owner/author(s). Publication rights licensed to ACM.
ACM ISBN 978-x-xxxx-xxxx-x/YYYY/MM
<https://doi.org/10.1145/nnnnnnnn.nnnnnnnn>

1 Introduction

Multilingual large language models (LLMs) have become essential infrastructure for global information access, achieving human-level performance across more than 50 languages. Models such as Llama-3.1 [7] and Qwen2.5 [11] demonstrate strong capabilities not only in English but also in Spanish, Chinese, Arabic, and other languages, suggesting the emergence of a **universal semantic space** shared across languages [5, 10]. However, a persistent gap remains: most models perform best in English, while accuracy in non-English languages is often 10–30% lower [9].

This gap is commonly attributed to the **English-centric imbalance in pretraining data**. English constitutes 70–80% of typical web-scale corpora [3], and limited exposure to other languages leads to performance degradation. Yet a fundamental question remains unanswered: *does the data distribution merely modulate performance, or does it fundamentally shape the geometry of the internal representation space?* If it is merely performance-dependent, post-hoc interventions such as fine-tuning or vocabulary adjustments may suffice; if it is structurally embedded, redesigning the pretraining stage is necessary.

Positioning within prior work. Prior studies have shown that multilingual models learn shared *cross-lingual spaces* [4, 10]. Early models such as mBERT [6] and XLM-R [5] demonstrated zero-shot transfer across languages, and more recent models such as Llama-3 and Qwen2.5 further strengthened cross-lingual capabilities with larger and more diverse corpora. Meanwhile, probing-based interpretability studies [1, 2, 8] examined how linguistic and syntactic information is encoded across layers, but did *not* systematically address (i) whether **linear separability of language information** is universal, (ii) at which depth **language directions emerge**, and (iii) how **pretraining language distributions are imprinted** into the model’s internal geometry. Our work fills this gap.

This study investigates two complementary research questions: (1) **Where and how is language information encoded within transformer layers?** (2) **Does the pretraining language distribution structurally reshape the geometry of multilingual representations?**

To answer these questions, we conduct a comprehensive probing study across *all 268 transformer layers* of six representative multilingual LLMs (Llama-3.1-8B, Qwen2.5-7B, Qwen2.5-Math-7B, OpenMath2-8B, OpenR1-7B, GPT-OSS-20B). We train linear and nonlinear probes and introduce a new analysis, **Token–Language Alignment**, across five typologically diverse languages (EN, ES, ZH, FR, DE), yielding a total of **2,680 independent experiments**.

Key Findings. (1) Universal Linear Separability. Across all models, language information is encoded in linearly accessible subspaces: linear probes achieve $99.8 \pm 0.1\%$ accuracy on average, with

only a $0.58 \pm 0.12\%$ gap relative to MLP probes ($p < 0.001$). Remarkably, this linear structure emerges *immediately at the first transformer block* (Layer 0 $\rightarrow$ 1), with a $+76.4 \pm 8.2\%$ jump, suggesting that language separation is an architecture-inherent mechanism.

(2) Structural Imprinting. Token–Language Alignment reveals that the alignment between language directions and vocabulary is strongly tied to the pretraining language distribution. English-centric models (EN > 80%) show 69.05% alignment for English but only 3.90% for Chinese, whereas Chinese-inclusive models (ZH20%) show a markedly higher 16.43% alignment for Chinese ($4.21\times$ increase). This indicates that the *geometry of representation space is structurally shaped by the pretraining distribution*, beyond its effects on task performance.

(3) Typological Effects and Complex Alignment Patterns. Chinese reaches its maximal separability in deeper layers (Layer 5.2 ± 0.8), whereas Spanish and German converge earlier (Layer 2.5 ± 0.4). The overall Match@Peak remains low ($15.0 \pm 0.3\%$), showing that language directions capture abstract, multi-factorial features beyond simple token–language matching, including script, lexical frequency, and typological structure.

Implications. These findings suggest that achieving fairness and balance in multilingual LLMs requires **careful design of pre-training data composition**. Structural imprinting implies that post-hoc adjustments cannot fundamentally alter the underlying geometry. Moreover, Match@Peak provides a quantitative diagnostic tool for evaluating the effects of data rebalancing on representation richness.

2 Methodology

We analyze how language information is encoded across all 268 transformer layers of six multilingual LLMs, and how such structure aligns with the pretraining language distribution. Our methodology consists of four parts: (1) experimental setup, (2) linear and non-linear probing, (3) Token–Language Alignment, and (4) evaluation and statistical analysis.

2.1 Experimental Setup

We study six multilingual LLMs whose pretraining corpora differ in their language composition:

- **English-centric models:** Llama-3.1-8B, OpenMath2-8B
- **Chinese-inclusive models:** Qwen2.5-7B, Qwen2.5-Math-7B, OpenR1-7B
- **Balanced baseline:** GPT-OSS-20B
- **Languages:** Following prior probing studies, we evaluate five XNLI languages (EN, ES, FR, DE, ZH), using 5k training and 2.5k validation sentences per language.

2.2 Probing Methodology

For each model and each layer ℓ , we use the hidden vector of the final token,

$$\mathbf{h}^{(\ell)} \in \mathbb{R}^d,$$

as the sentence representation.

Linear probe. A linear classifier is applied to LayerNorm-normalized representations:

$$f_{\text{lin}}(\mathbf{h}) = W_c \cdot \text{LN}(\mathbf{h}) + b_c,$$

trained with cross-entropy loss to predict one of the five languages.

MLP probe. To compare linear and nonlinear capacity, we additionally train an MLP:

$$f_{\text{mlp}}(\mathbf{h}) = W_2 \cdot \text{ReLU}(W_1 \cdot \text{LN}(\mathbf{h})).$$

LayerNorm is applied in both probes to remove inter-layer scale differences and measure *pure linear separability* of language information.

Both probes are trained independently for each (model, layer, seed, probe type) using AdamW ($\text{lr} = 10^{-3}$), batch size 128, and 3 epochs with early stopping.

2.3 Token–Language Alignment

To quantify how learned language directions relate to the LM head vocabulary, we compute cosine similarity between each probe-learned language direction $\mathbf{w}_L^{(\ell)}$ and each vocabulary embedding $\mathbf{e}_v$:

$$\text{sim}(v, L, \ell) = \cos(\mathbf{e}_v, \mathbf{w}_L^{(\ell)}).$$

Each token v is assigned to the language direction with highest similarity.

For each language L , we compute three alignment metrics:

- **PeakDepth_L:** The normalized layer index where $\text{VocabShare}(\ell, L)$ is maximized, indicating the depth at which language L is most strongly expressed.
- **PeakVocab_L:** The maximum value of $\text{VocabShare}(\ell, L)$ across layers, representing how dominant language L is within the vocabulary space.
- **Match@Peak_L:** Among tokens assigned to language L at its peak layer, the percentage whose decoded text belongs to language L , determined using Unicode and diacritic rules (e.g., CJK blocks for ZH, accent-sensitive characters for ES/FR/DE). Higher values indicate stronger alignment between the learned direction and the true lexical identity.

Together, these metrics measure (i) where language information appears in the network, (ii) how strongly it organizes the vocabulary, and (iii) how closely the learned directions correspond to actual linguistic identity—capturing the *structural imprinting* left by pretraining data.

2.4 Evaluation and Statistics

Probe accuracy is evaluated on the validation set. Differences between linear and MLP probes are assessed using layer-wise paired t-tests. We report mean accuracy and 95% confidence intervals across multiple random seeds to ensure robustness and statistical reliability.

3 Experiment Result

3.1 Layer-wise Language Separability

As shown in Table 1, all transformer layers except the initial embedding layer (Layer 0) achieve consistently high language classification accuracy, exceeding 90% across all models. This indicates that language information is not a transient or layer-specific signal but rather a **global representational property that is preserved throughout model depth**. The sharp increase in accuracy from

Table 1: Layer-wise Performance Statistics for 6 Multilingual Models Across 5 Languages (with per-model language average). Values show mean accuracy (%) with standard deviation where applicable. Gap significance: *p<0.05, **p<0.01, *p<0.001. Right block adds Token–Language Alignment: PeakDepth (normalized depth of vocab-share peak), PeakVocab (%), Match@Peak (%).**

Model	Lang	Early Layer Dynamics			Linear Separability			Aggregate		Token–Language Alignment		
		L0	L1	Jump	Linear Avg	MLP Avg	Gap	Last	Avg	PeakDepth	PeakVocab	Match@Peak
Llama-3.1-8B	EN	20.0	99.8	79.8	97.3±1.5	97.9±1.8	+0.6	100.0	97.3	0.06	39.2	67.9
	ES	20.0	99.8	79.8	96.5±1.7	95.5±0.8	-1.0	99.8	96.5	0.71	32.1	1.8
	ZH	20.0	99.5	79.5	96.0±1.9	96.1±2.3	+0.1	100.0	96.0	0.13	38.9	4.1
	FR	20.0	99.7	79.7	94.9±2.1	95.2±3.6	+0.3	99.6	94.9	0.06	26.9	0.8
	DE	20.0	100.0	80.0	96.2±1.8	95.3±0.6	-0.9	99.8	96.2	1.00	28.6	0.4
	Avg	20.0	99.8	79.8	96.2±1.8	96.0±1.8	-0.2	99.8	96.2	0.39	33.2	15.0
Qwen2.5-7B	EN	38.3	99.4	61.1	97.4±2.0	98.5±2.0	+1.1	98.6	97.4	0.59	40.8	52.6
	ES	16.7	99.6	82.9	94.6±2.2	93.8±3.0	-0.8	98.7	94.6	0.15	35.3	0.8
	ZH	2.4	94.7	92.3	94.3±0.5	93.9±1.1	-0.3	99.6	94.3	0.44	23.1	17.7
	FR	32.0	99.6	67.6	95.7±2.3	93.7±1.5	-1.9	98.3	95.7	0.04	33.9	0.8
	DE	0.0	99.7	99.7	94.3±0.6	93.9±1.1	-0.3	98.9	94.3	0.19	29.3	0.3
	Avg	17.9	98.6	80.7	95.2±1.5	94.8±1.7	-0.5	98.8	95.2	0.28	32.5	14.4
Qwen2.5-Math-7B	EN	38.3	99.8	61.5	97.5±2.0	98.6±1.9	+1.1	99.9	97.5	0.30	51.8	53.9
	ES	16.7	99.5	82.8	95.0±2.0	94.2±3.0	-0.8	98.8	95.0	0.00	28.3	1.0
	ZH	2.4	76.2	73.8	80.8±4.3	83.5±6.1	+2.7	85.4	80.8	0.81	57.4	15.7
	FR	16.0	99.8	83.8	95.1±2.0	93.7±1.9	-1.4	99.3	95.1	0.96	46.0	0.8
	DE	15.7	99.5	83.9	96.0±1.6	95.0±0.9	-1.1	99.7	96.0	1.00	57.6	0.3
	Avg	17.8	95.0	77.1	92.9±2.4	93.0±2.8	+0.1	96.6	92.9	0.61	48.2	14.3
OpenMath2-8B	EN	20.0	99.8	79.8	96.9±1.6	97.2±2.4	+0.4	98.8	96.9	0.55	25.9	70.2
	ES	0.0	99.4	99.4	95.1±0.6	94.9±1.0	-0.2	98.6	95.1	0.68	22.5	1.4
	ZH	0.0	99.4	99.4	92.9±1.5	93.8±3.3	+0.9	94.5	92.9	0.23	53.9	3.7
	FR	40.0	99.8	59.8	95.8±2.4	96.0±2.7	+0.2	97.8	95.8	0.03	39.4	0.9
	DE	40.0	100.0	60.0	97.1±2.1	96.0±0.6	-1.1	99.4	97.1	1.00	32.3	0.4
	Avg	20.0	99.7	79.7	95.5±1.6	95.6±2.0	+0.0	97.8	95.5	0.50	34.8	15.3
OpenR1-7B	EN	57.4	99.7	42.3	98.1±2.0	98.3±2.2	+0.1	99.5	98.1	0.00	60.6	55.9
	ES	16.7	99.8	83.1	96.5±1.6	96.1±2.1	-0.4	99.4	96.5	0.11	24.4	0.9
	ZH	2.4	65.0	62.6	83.9±2.3	85.8±5.0	+2.0	88.7	83.9	0.52	67.4	15.9
	FR	16.0	99.9	83.9	95.6±1.9	94.1±2.0	-1.5	99.1	95.6	0.59	64.7	0.8
	DE	0.0	99.7	99.7	95.9±0.2	95.2±1.1	-0.7	99.5	95.9	0.04	25.0	0.3
	Avg	18.5	92.8	74.3	94.0±1.6	93.9±2.5	-0.1	97.2	94.0	0.25	48.4	14.8
GPT-OSS-20B	EN	38.8	99.2	60.4	96.5±2.5	99.2±0.3	+2.7	98.0	96.5	0.74	35.1	65.0
	ES	34.6	90.6	56.0	88.1±5.9	89.8±2.8	+1.7	96.9	88.1	0.61	23.3	2.0
	ZH	26.2	96.9	70.8	71.6±6.7	80.1±6.0	+8.5	59.2	71.6	0.13	31.0	6.0
	FR	0.0	90.7	90.7	84.9±3.9	90.2±3.1	+5.3	93.0	84.9	0.52	28.0	1.0
	DE	0.0	89.4	89.4	86.2±2.5	90.2±3.3	+4.1	94.8	86.2	0.00	23.2	0.8
	Avg	19.9	93.4	73.5	85.5±4.3	89.9±3.1	+4.5	88.4	85.5	0.40	28.1	14.9

Table 2: Match@Peak (%) by model group, grouped by pre-training language distribution.

Model Group	EN	ZH	ES	FR	DE
English-centric ^a	69.05	3.90	1.60	0.85	0.40
Chinese-inclusive ^b	54.13	16.43	0.90	0.80	0.30
Difference ($\Delta\%$)	14.92	12.53	0.70	0.05	0.10
Ratio ($\times$)	1.28	4.21	1.78	1.06	1.33

^aLlama-3.1-8B, OpenMath2-8B: models with English-centric pretraining (ZH share < 5%).

^bQwen2.5-7B, Qwen2.5-Math-7B, OpenR1-Qwen-7B: models trained on both English and Chinese (ZH-inclusive pretraining).

Layer 0 to Layer 1 (+76.4±8.2%p) suggests the presence of an **early structural reorganization stage** in which language information is rapidly separated within the first transformer block.

Furthermore, the performance gap between linear and MLP probes is less than 1%p on average, demonstrating that language information does not require nonlinear decision boundaries. Instead, it is **almost fully linearly separable** within the latent space. This implies that language is encoded not as a complex nonlinear pattern but as a **global directionality** in a high-dimensional representation space.

These observations point to two structural properties. First, the model appears to **establish a normalized representation of language identity at very early layers** and preserve this signal throughout the stack. Second, the strong linear separability supports the hypothesis that **languages form low-dimensional yet coherent subspaces** in the latent representation space. These findings are consistent with the Token–Language Alignment and structural imprinting analyses presented in the next subsection.

In summary, our layer-wise probing experiments show that **LLM latent spaces preserve language information in a clearly and**

structurally separable form, even without additional nonlinear capacity.

3.2 Token–Language Alignment and Structural Imprinting

Token–Language Alignment evaluates how language directions learned by the linear probe relate to LM head token embeddings. Our results show that **LLMs do not distinguish languages solely based on Unicode-defined language identity**. In particular, the alignment between language directions and token identity (Match@Peak) is very low for Latin-script languages such as Spanish, French, and German. This suggests that the learned language directions reflect **latent representational structures shaped during pretraining**, rather than surface-level script information.

In contrast, Chinese (ZH) exhibits substantially higher alignment. Chinese-inclusive models (with $\approx 20\%$ ZH data) reach a mean Match@Peak of 16.43%, whereas English-centric models (EN $>80\%$) reach only 3.90%, indicating a $4.21\times$ increase. This demonstrates that the **pretraining language distribution can directly influence the geometry of latent representations**, forming clearer language-specific directions when a language is sufficiently represented in the training corpus.

Taken together, these observations suggest the following structural characteristics of multilingual LLMs:

- (1) **Rather than relying on surface-level language ID, the model organizes languages via latent directions in the representation space that emerge during pretraining.** These directions correspond to low-dimensional classification boundaries learned by the probe and capture differences between language-specific representations.
- (2) **Languages with sufficient pretraining coverage (e.g., ZH) form clearer and more stable latent directions**, whereas languages with limited data or shared scripts (ES/FR/DE) tend to share mixed and less distinct subspaces.

Overall, our analysis demonstrates that **LLMs separate languages not by surface token-level features but by latent representational structures shaped by the pretraining distribution**. Moreover, **the composition of pretraining data plays a decisive role in determining the geometry of multilingual representations**.

4 Conclusion

This work presents a quantitative analysis of how language information is structured across all 268 transformer layers of six multilingual LLMs. Our experiments show that language separation emerges immediately in the first transformer block and remains a **stable and strongly linear structure** throughout model depth. The negligible performance gap between linear and MLP probes indicates that the information required for distinguishing languages resides in a **linearly accessible latent subspace**, rather than in complex nonlinear boundaries.

Through our proposed Token–Language Alignment analysis, we further observe that the alignment between language directions and vocabulary embeddings is **strongly influenced by the language composition of the pretraining data**. Chinese (ZH) exhibits much clearer alignment in models that include substantial

ZH data, whereas languages with limited coverage or shared scripts (ES/FR/DE) show weaker alignment. These findings demonstrate that **the pretraining corpus leaves a structural imprint on the geometry of multilingual representations**.

Overall, our results indicate that language representation in LLMs is distinguished not by surface-level token features, but by **latent representational structures formed during pretraining**, which are established early and preserved consistently across layers. This highlights the importance of **balanced language composition and corpus design** in achieving fairness and robustness in multilingual LLMs. Moreover, Match@Peak and token–language alignment metrics provide practical tools for diagnosing representation richness and distributional biases.

Future directions include expanding the analysis to additional languages and scripts, evaluating the influence of tokenizer design, and conducting interventional studies to further investigate the causal role of language directions in model behavior.

References

- [1] Guillaume Alain and Yoshua Bengio. 2016. Understanding Intermediate Layers Using Linear Classifier Probes. *arXiv preprint arXiv:1610.01644* (2016). doi:10.48550/arXiv.1610.01644
- [2] Yonatan Belinkov. 2022. Probing Classifiers: Promises, Shortcomings, and Advances. *Computational Linguistics* 48, 1 (2022), 207–219. doi:10.1162/coli_a_00422
- [3] Emily M. Bender, Timnit Gebru, Angelina McMillan-Major, and Shmargaret Shmitchell. 2021. On the Dangers of Stochastic Parrots: Can Language Models Be Too Big?. In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency (FAccT '21)*. Association for Computing Machinery, New York, NY, USA, 610–623. doi:10.1145/3442188.3445922
- [4] Ethan A. Chi, John Hewitt, and Christopher D. Manning. 2020. Finding Universal Grammatical Relations in Multilingual BERT. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*. Association for Computational Linguistics, Online, 5564–5577. doi:10.18653/v1/2020.acl-main.493
- [5] Alexis Conneau, Kartikay Khandelwal, Naman Goyal, Vishrav Chaudhary, Guillaume Wenzek, Francisco Guzmán, Edouard Grave, Myle Ott, Luke Zettlemoyer, and Veselin Stoyanov. 2020. Unsupervised Cross-lingual Representation Learning at Scale. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*. Association for Computational Linguistics, Online, 8440–8451. doi:10.18653/v1/2020.acl-main.747
- [6] Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*. Association for Computational Linguistics, Minneapolis, Minnesota, 4171–4186. doi:10.18653/v1/N19-1423
- [7] Aaron Grattafiori, Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, et al. 2024. The Llama 3 Herd of Models. *arXiv preprint arXiv:2407.21783* (2024). <https://arxiv.org/abs/2407.21783>
- [8] John Hewitt and Christopher D. Manning. 2019. A Structural Probe for Finding Syntax in Word Representations. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*. Association for Computational Linguistics, Minneapolis, Minnesota, 4129–4138. doi:10.18653/v1/N19-1419
- [9] Junjie Hu, Sebastian Ruder, Aditya Siddhant, Graham Neubig, Orhan Firat, and Melvin Johnson. 2020. XTREME: A Massively Multilingual Multi-task Benchmark for Evaluating Cross-lingual Generalization. In *Proceedings of the 37th International Conference on Machine Learning*. arXiv:2003.11080 <https://arxiv.org/abs/2003.11080> ICMEL.
- [10] Telmo Pires, Eva Schlinger, and Dan Garrette. 2019. How Multilingual is Multilingual BERT?. In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*. Association for Computational Linguistics, Florence, Italy, 4996–5001. doi:10.18653/v1/P19-1493
- [11] An Yang, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, et al. 2024. Qwen2.5 Technical Report. *arXiv preprint arXiv:2412.15115* (2024). doi:10.48550/arXiv.2412.15115