



VITA TOTA DISCENDA EST

Transformers: You Need Attention!

Name: JIAHUI ZHUO GAO

Institute: IFIC (UV - CSIC)

Date: 21/09/2026

Raymundus Lullus training program in advanced computing



UNIVERSIDADE DA CORUÑA



Universidad de Oviedo



URL



Centro de Investigaciones
Energéticas, Medioambientales
y Tecnológicas



IPARCOS

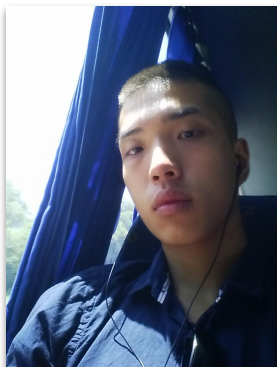


Artemisa

About the lecturer



VITA TOTA DISCENDA EST



- **JIAHUI ZHUO GAO**
- **Ph.D., Universitat de València (IFIC-CSIC), 2026**
- **LHCb Trigger & Real-Time Analysis**
- **One of main developer & maintainer, LHCb HLT1 (Run 3)**
- **Convener, LHCb Retina DWT Simulation & Software WG**
- **2024 LHCb Early Career Scientist Award**

Raymundus Lullus training program in advanced computing



UNIVERSIDADE DA CORUÑA



Universidad de Oviedo





VITA TOTA DISCENDA EST

Transformers: You Need Attention!

35' of fundamentals

- What is transformer?
- Does it have intelligence?
- How could we use it?

45' hands-on

- Demonstrate the intelligence
- Example of usage

Raymundus Lullus training program in advanced computing





VITA TOTA DISCENDA EST

What is transformer?

Raymundus Lullus training program in advanced computing



VITA TOTA DISCENDA EST



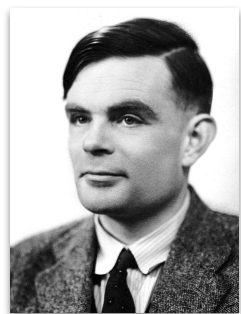
Introduction

The iPhone 4 moment



Alan Mathison Turing

“Computing Machinery and Intelligence”



Transformer

“Attention Is All You Need”



ChatGPT

1M users in 5 days
100M in 2 months



Raymundus Lullus training program in advanced computing





VITA TOTA DISCENDA EST

Transformer

The cat sat on the mat because it was



Transformer



The cat sat on the mat because it was **tired**

Raymundus Lullus training program in advanced computing





VITA TOTA DISCENDA EST

Transformer

The cat sat on the mat because it was



Transformer

How is this possible?



The cat sat on the mat because it was **tired**

Raymundus Lullus training program in advanced computing



UNIVERSIDADE DA CORUÑA



UNIVERSIDADE DE SANTIAGO DE COMPOSTELA



Universidad de Oviedo



IFCA



URL

IFIC



Ciemat
Centro de Investigaciones
Energéticas, Medioambientales
y Tecnológicas

UAM



IPARCOS

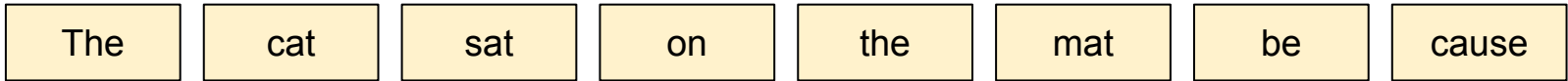
CAPA



Artemisa

Text becomes a row of vectors

tokenize the text



embed the token into vector



Raymundus Lullus training program in advanced computing



VITA TOTA DISCENDA EST

Attention head

The	i-3
cat	i-2
sat	i-1
?	i

- W_Q, W_K, W_V are weights of attention head
- d_h is the dimension of attention head
- W_O is the transformation matrix

$$\mathbf{t}_i \leftarrow \mathbf{t}_i + \sum_{j=0}^i \text{softmax}_j \left[\frac{(\mathbf{t}_i W_Q)(\mathbf{t}_j W_K)^{\top}}{\sqrt{d_h}} \right] (\mathbf{t}_j W_V) W_O$$

Q: query

K: key

V: value

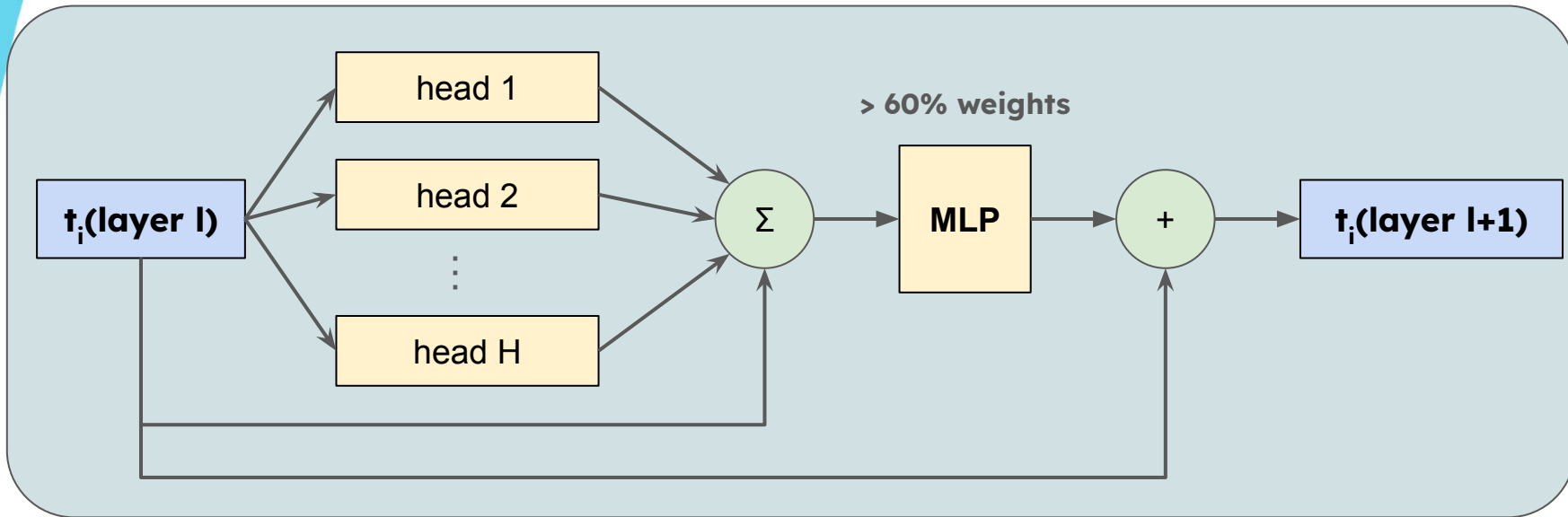
Raymundus Lullus training program in advanced computing





VITA TOTA DISCENDA EST

Attention layer



In high dimensions almost all directions are nearly orthogonal.

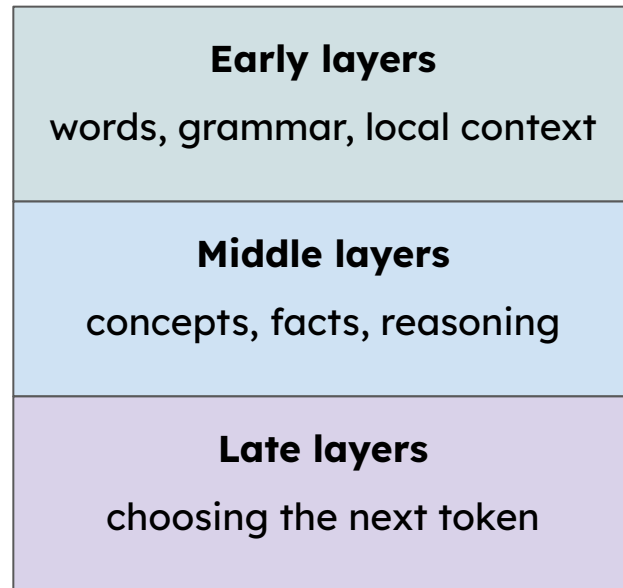
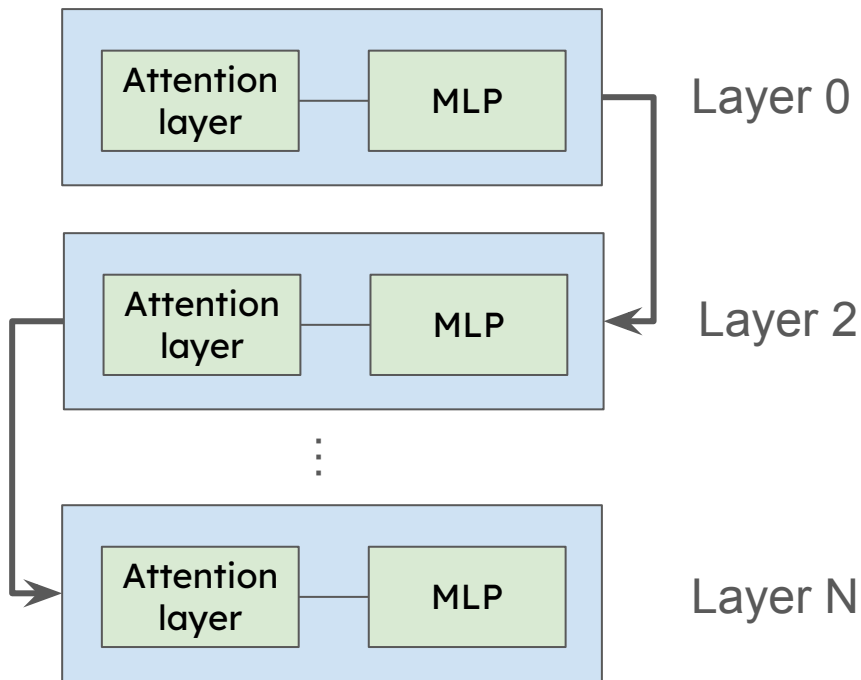
Raymundus Lullus training program in advanced computing





VITA TOTA DISCENDA EST

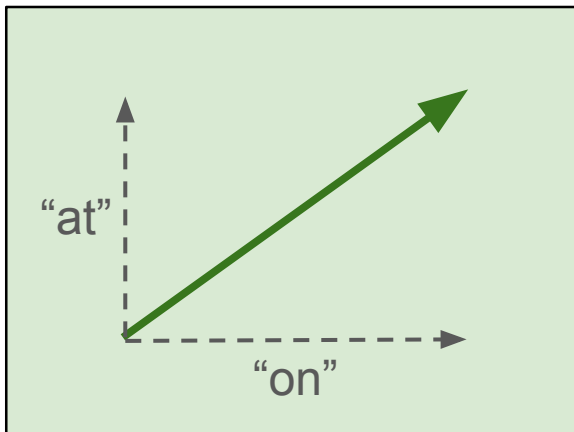
Transformer layers



Raymundus Lullus training program in advanced computing



Transformer output



Its projection onto each vocabulary token gives that token's probability of coming next (after normalization).

The **output vector** of the last layer is the prediction for the next token.

Raymundus Lullus training program in advanced computing





VITA TOTA DISCENDA EST

Does it have intelligence?

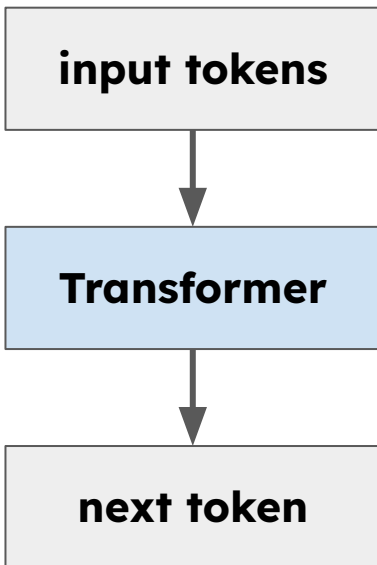
Raymundus Lullus training program in advanced computing





VITA TOTA DISCENDA EST

Transformer



That is all it does.
Predict the next token



Is this intelligence?

Or just a very good
guess at the next token?

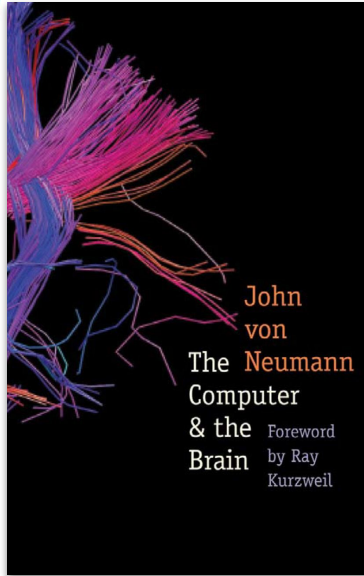
Raymundus Lullus training program in advanced computing



John von Neumann, 1958



VITA TOTA DISCENDA EST



“...precisions like the ones mentioned above (10 to 12 decimals!) are altogether out of the question. The nervous system is a computing machine which manages to do its exceedingly complicated work at a rather low level of precision (2 to 3 decimals) ... no known computing machine can operate reliably and significantly at such a low precision level.

... [this] leads not only to a low level of precision but also to a rather high level of reliability ... if ... a single pulse is lost, or even several pulses are lost ... the relevant frequency, i.e. the meaning of the message, is only inessentially distorted.”

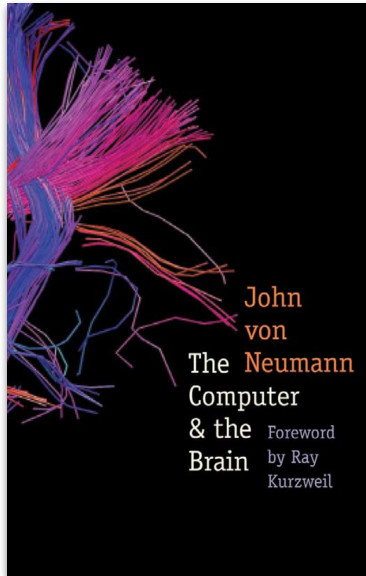
Raymundus Lullus training program in advanced computing



John von Neumann, 1958



VITA TOTA DISCENDA EST



"The nervous system appears to be using a radically different system of notation from the ones we are familiar with in ordinary arithmetics and mathematics ... the meaning is conveyed by the statistical properties of the message ... a deterioration in arithmetics has been traded for an improvement in logics."

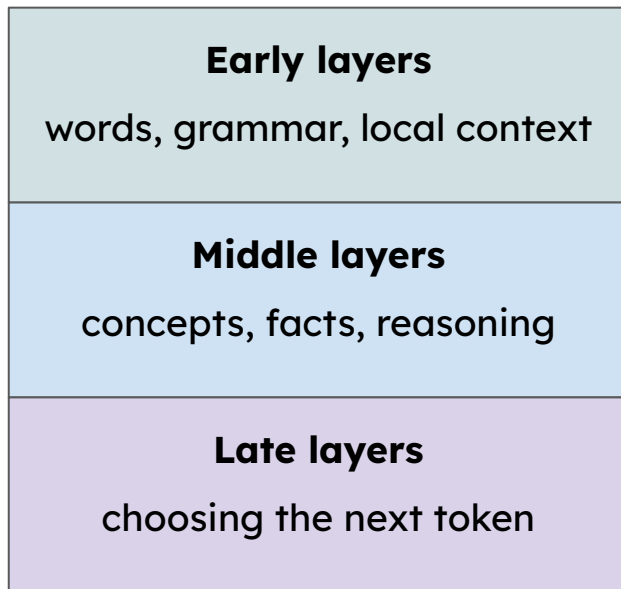
The brain's notation is statistical

So is the Transformer's.

Raymundus Lullus training program in advanced computing



So where is the understanding?



If transformer has understanding, it has to be written somewhere.

Can we read it?

Can we edit it?

Raymundus Lullus training program in advanced computing





VITA TOTA DISCENDA EST

j-space

$$J_\ell = \mathbb{E}_{t, t' \geq t, \text{prompt}} \left[\frac{\partial h_{\text{final}, t'}}{\partial h_{\ell, t}} \right], \quad \text{lens}(h_\ell) = \text{softmax}(W_U \text{norm}(J_\ell h_\ell))$$

Verbalizable Representations Form a Global Workspace in Language Models

Wei Gurneus*, Nicholas Sotgiu*, Adam Pearce, Mateusz Piotrowski, Isaac Kuvshinov, Ranjin Chen, Anna Soltes, Paul Bhagavat, Eason Ong, Rowan Wang, Ben Thompson, David Abulafia, Subhash Kantamannil, Emmanuel Ameisen, Joshua Batson

Jack Linney¹

* Core contributor; ¹ Correspondence to jacklinney@anthropic.com

Abstract

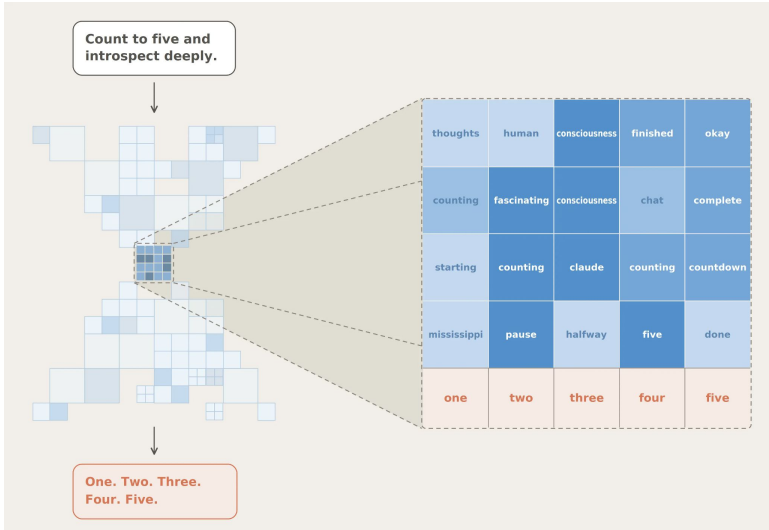
Out of everything the human brain processes, only a small fraction is consciously accessible, in the sense of being available for verbal report, deliberate control, and flexible reasoning. In this paper, we present evidence that an analogous functional distinction has emerged in large language models. Using a new interpretability technique, the Jacobian lens, we identify the representations a model is poised to verbalize at any point in its processing. These representations, which we collectively call the J-space, exhibit the functional properties characteristic of a global workspace: their contents can be reported, deliberately summoned and held, used to carry the intermediate steps of silent reasoning, and passed as arguments to arbitrary downstream computations, while automatically processing such as text parsing and routine inference proceeds without them. The J-space also has structural signatures that global workspace theory associates with conscious access: it carries coherent content only in an intermediate band of layers, holds on the order of tens of concepts at a time, and is broadcast by the model's weights more widely than other representations. These properties make it a practical window into a model's implicit thinking. In alignment audits, it reveals strategic deliberation, evaluation awareness, and trained-in misaligned dispositions that never appear in the model's outputs. We find that post-training installs the Assistant's point of view in the workspace, and we introduce counterfactual reflection training, which improves behavior by training only what a model would say if interpreted and asked to reflect. These results indicate that language models maintain a small, privileged set of representations bearing some of the functional hallmarks of conscious access, and that decoding these representations sheds light on ongoing cognitive processes.

1 Introduction

If the mind is an ocean, we spend our lives floating at the surface. Beneath us, an enormous amount of processing takes place without our knowledge: our visual systems parsing the contours of a face, our motor circuits maintaining our posture. At any given moment, only a small fraction of this neural activity is accessible to us. Yet it is this privileged sliver of activity that we rely on to reason deliberately: to plan what ingredients to buy for a recipe, or to puzzle out why an engine won't start. Such thoughts can be articulated out loud, deliberately held in mind, and brought to bear on whatever task the moment demands. This distinction, between our accessible thoughts and our unconscious processing, is perhaps the most striking feature of human cognition.

Archival Preprint. We recommend reading the [HTML](#) render for improved styling and interactivity.

arXiv:2607.15495v1 [cs.CL] 16 Jul 2026



Raymundus Lullus training program in advanced computing



Functional roles of j-space

Verbal report

What are you thinking about?

J-space

banana

elephant

"...about a
banana."
"...about an
elephant."

Directed modulation

Compute $3^2 - 2$ while
writing "The old painting
hung crookedly on the
wall."

J-space

math - calc - nine
- seven = equals

"The old
painting
hung
crookedly
on the
wall."

Internal reasoning

What color is the planet
fourth from the sun?

J-space

Mars

Earth

"red"
"blue"

Raymundus Lullus training program in advanced computing



Functional roles of j-space

Flexible generalization

What is the capital of France?

J-space

France

China

capital → Paris	Beijing
language → French	Chinese
continent → Europe	Asia
currency → Euro	Yuan

Selectivity

(a battery of tasks)

J-space

~~task plan step~~
~~reason answer~~

- ✓ parse inputs
- ✓ recall facts
- ✓ speak fluently
- ✗ reason internally
- ✗ make complex inferences

With the J-lens we can watch the model think.

Raymundus Lullus training program in advanced computing





VITA TOTA DISCENDA EST

How could we use it?

Raymundus Lullus training program in advanced computing





Vision Banana

Google DeepMind

2026-06-05

Image Generators are Generalist Vision Learners

Valentin Gabeur¹, Shangbang Long², Songze Peng³, Paul Voigtlaender⁴, Shiyang Sun⁵, Yanan Bao⁶, Karen Tsung⁷, Zheheng Wang⁸, Weslei Zhou⁹, Jonathan T. Barron¹, Kyle Genova¹, Nikhith Kannan¹, Sherry Ben¹, Yandong Li¹, Mandy Guo¹, Sahar Vega¹, Yanning Gu¹, Huizhong Chen¹, Oliver Wang¹, Saining Xie¹, Howard Zhou¹, Kaiming He¹, Thomas Funkhouser¹, Jean-Baptiste Ayoub¹ and Kato Saito¹

¹project leads and equal contribution, ²core contributors, ³project advisors, ⁴leadership sponsors

Recent works show that image and video generators exhibit zero-shot visual understanding behaviors, in a way reminiscent of how Large Language Models (LLMs) develop emergent capabilities of language understanding and reasoning from generative pretraining. While it has long been conjectured that the ability to create visual content implies an ability to understand it, there has been limited work demonstrating that generalist image generators can achieve state-of-the-art understanding capabilities on many different vision tasks. In this work, we demonstrate that image generation training serves a role similar to LLM pretraining, and lets models learn powerful and general visual representations that enable state-of-the-art performance on various vision tasks. We introduce *Vision Banana*, a generalist model built by instruction-tuning *Nano Banana Pro* (NBP) on a mixture of its original training data alongside a small amount of vision task data. By parameterizing the output space of vision tasks as RGB images, we seamlessly reframe perception as image generation. Our generalist model, *Vision Banana*, achieves state-of-the-art results on a variety of vision tasks involving both 2D and 3D understanding, beating or rivaling zero-shot domain-specialists, including *Segment Anything Model 3* on segmentation tasks, and the *Depth Anything* series on metric depth estimation. We show that these results can be achieved with lightweight instruction-tuning without sacrificing the base model's image generation capabilities. The superior results suggest that image generation pretraining is a generalist vision learner. It also shows that image generation serves as a unified and universal interface for vision tasks, similar to text generation's role in language understanding and reasoning. We could be witnessing a major paradigm shift for computer vision, where generative vision pretraining takes a central role in building Foundational Vision Models for both generation and understanding.

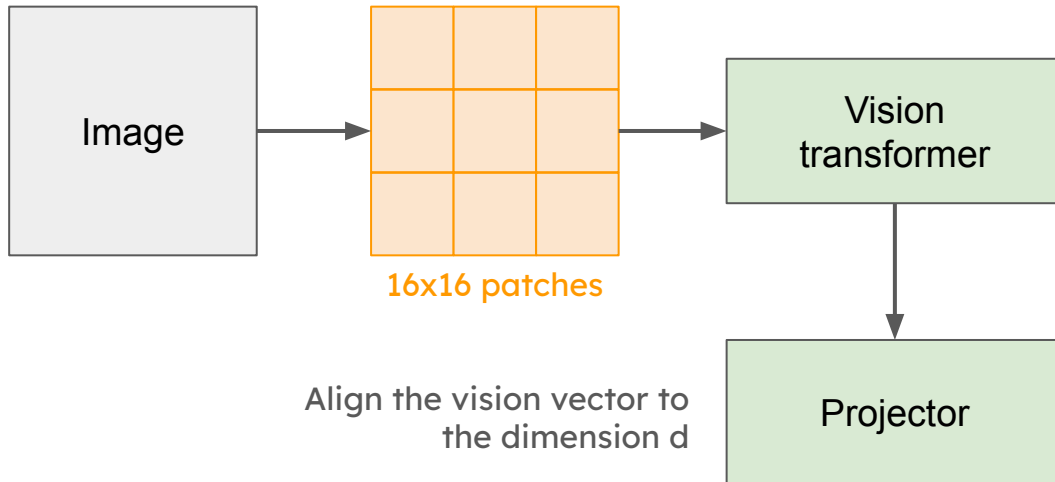
Project Page: vision-banana.github.io

1. Introduction

In recent years, advanced image and video generation models (Black Forest Labs, 2025; ByteDance, 2026; Google, 2025a,b; Luma, 2026; OpenAI, 2026) have demonstrated unprecedented generation capabilities, synthesizing highly complex, high-fidelity visual content with precise semantic control. This remarkable capability for visual creation suggests that these models possess a deep, internalized comprehension of the visual world's underlying structures, semantics, and relationships. However, leading methods on visual representation learning in general do not belong to the family of generative modeling. Instead, they include supervised discriminative learning (Dehghani et al., 2023; Dosovitskiy et al., 2020; Krizhevsky et al., 2012), contrastive learning (Chen et al., 2020b,c; He et al., 2020; Radford et al., 2021; Tschannen et al., 2025; Zhai et al., 2023), bootstrapping (Caron et al., 2021; Grill et al., 2020), auto-encoding (Bao et al., 2021; Chen et al., 2024; He et al., 2022) among others, and their combinations (Cao et al., 2026; Oquab et al., 2023; Siméoni et al., 2025; Zhou et al., 2021). Early efforts in generative vision pretraining (Bai et al., 2024; Chen et al., 2020a) have shown promising scaling behaviors but their effectiveness has lagged behind non-generative models.

Correspondence: vision-banana@google.com
© 2026 Google. All rights reserved.

How images enter the transformer?



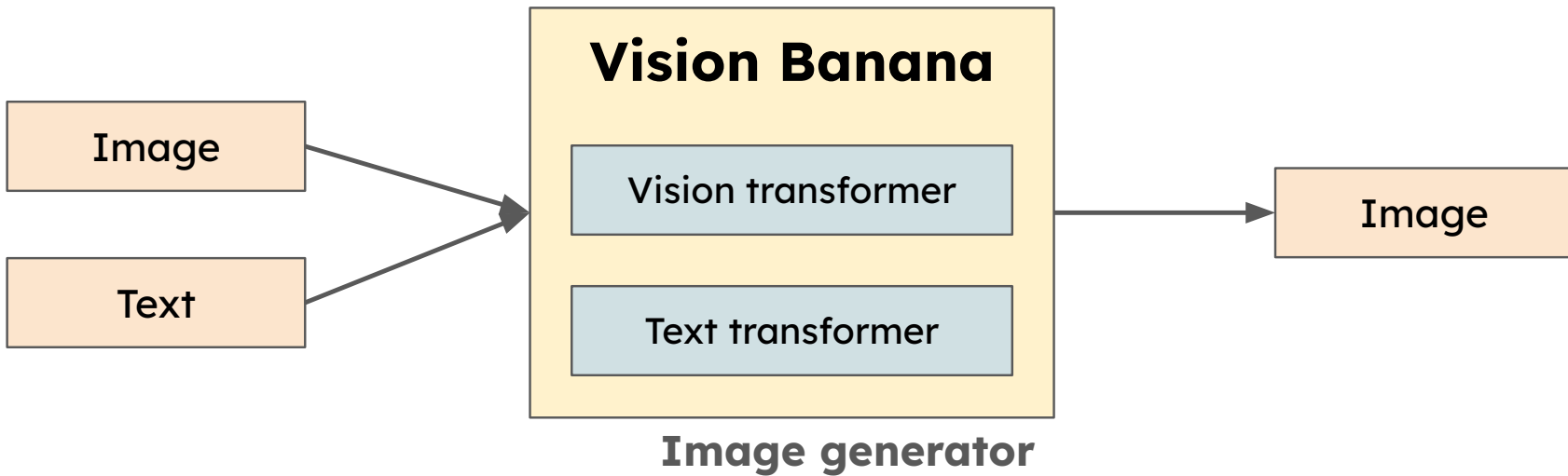
Raymundus Lullus training program in advanced computing





VITA TOTA DISCENDA EST

Vision Banana

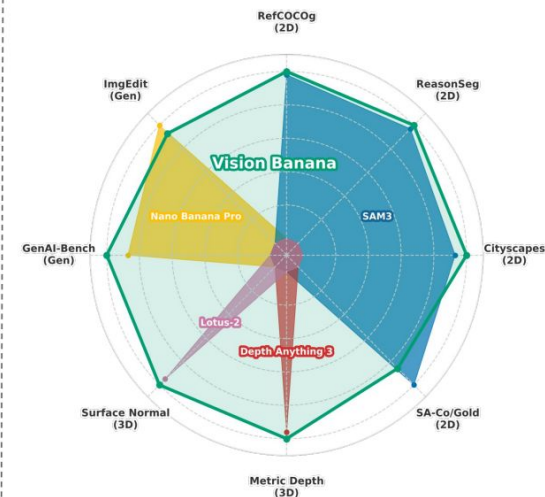
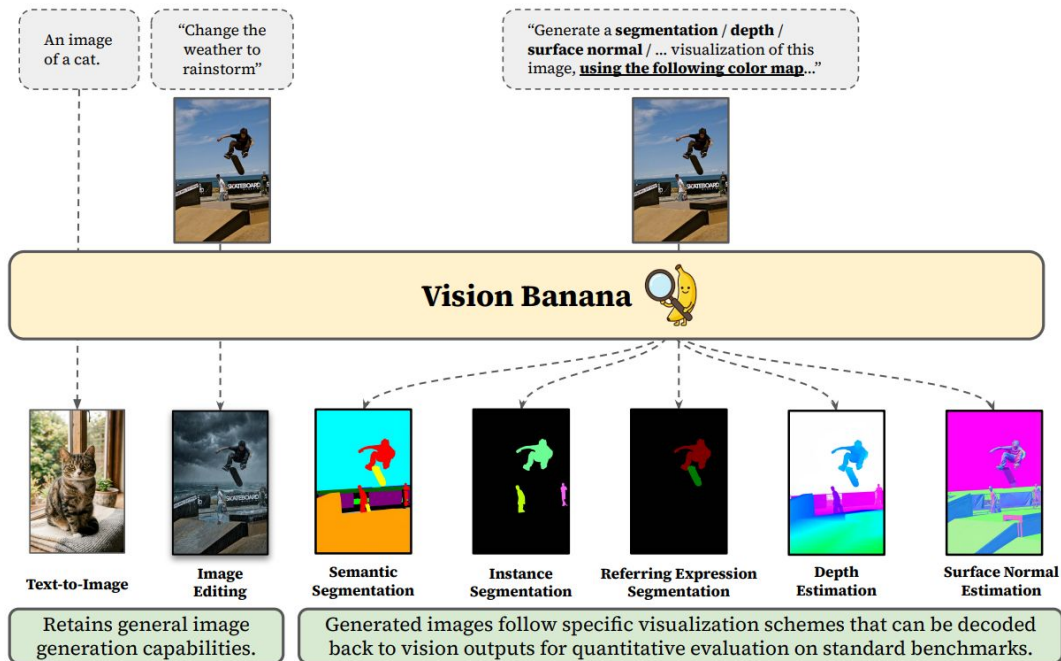


Raymundus Lullus training program in advanced computing





Vision Banana



Vision Banana is a generalist vision model in both visual generation and understanding, surpassing or rivaling specialist models.

Raymundus Lullus training program in advanced computing



Give Transformer new knowledge

In the context

No setup

Limited by context size

Attention gets diluted in long inputs

Fine tuning

Knowledge goes into the weights

Expensive, retrain on every update

Good for style and format, unreliable for facts

RAG

Retrieval-Augmented Generation

Retrieve the relevant information, put them in the context

Scales to large document sets, cheap to update

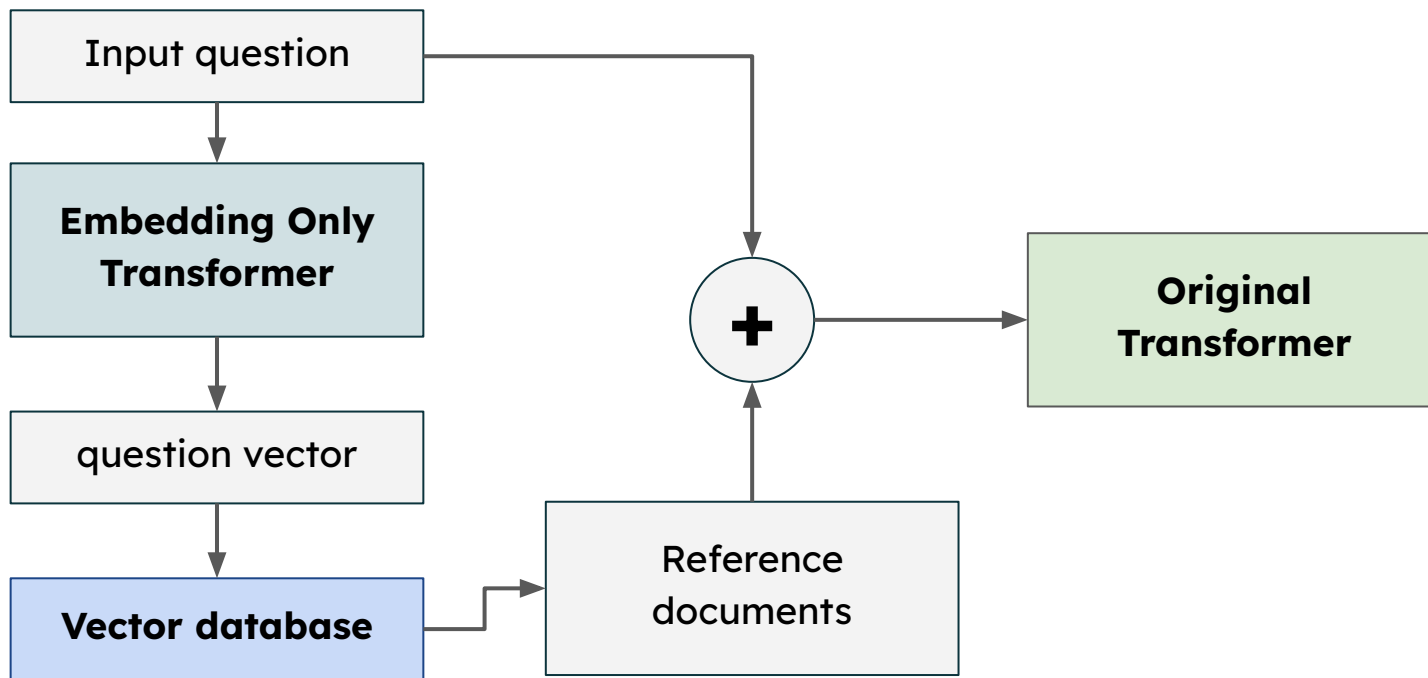
Raymundus Lullus training program in advanced computing





VITA TOTA DISCENDA EST

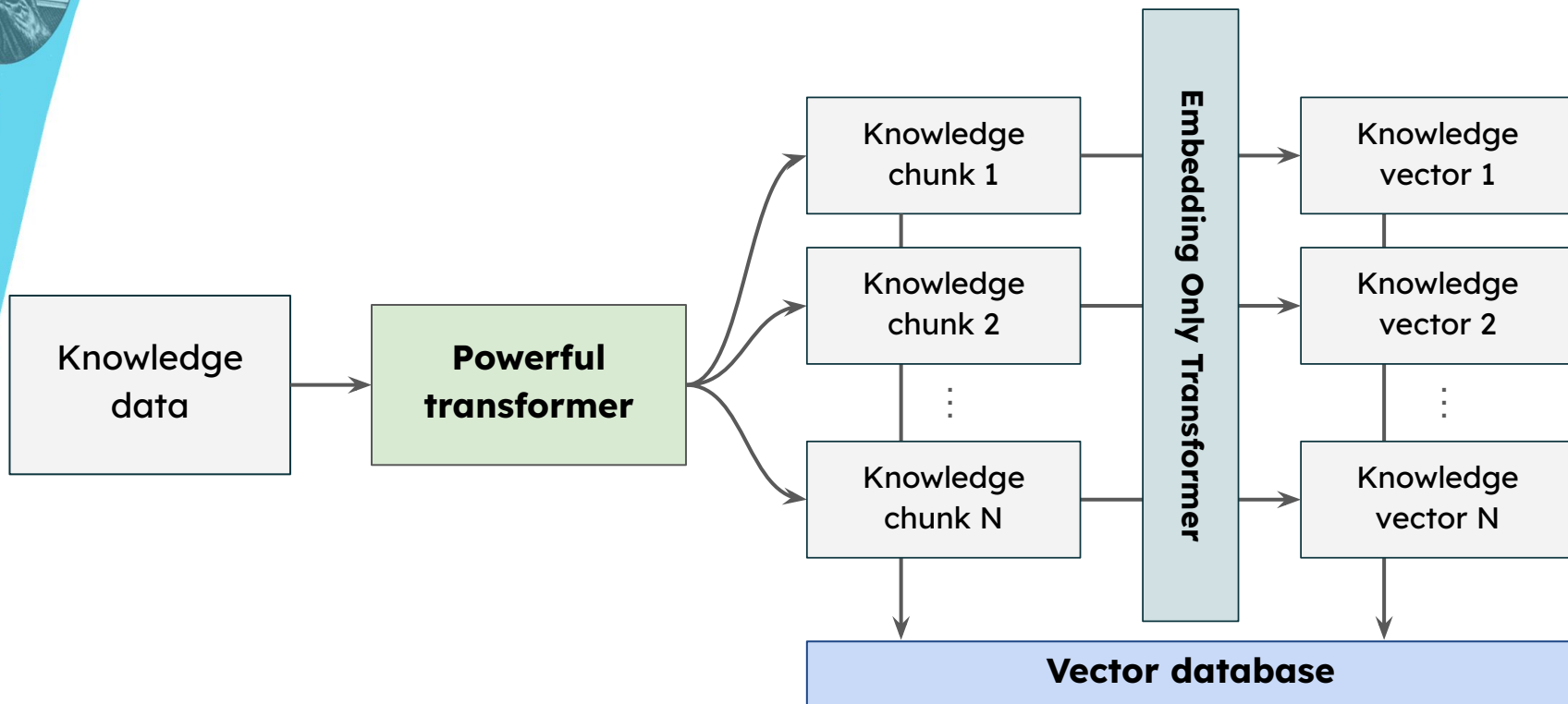
Retrieval-Augmented Generation



Raymundus Lullus training program in advanced computing



Retrieval-Augmented Generation



Raymundus Lullus training program in advanced computing





Transformer in HEP

arXiv:2412.05863v1 [hep-ex] 8 Dec 2024

Run 3 performance and advances in heavy-flavor jet tagging in CMS

Uttiya Sarkar on behalf of the CMS collaboration^{a,*}
^aIII. Physikalisches Institut A, RWTH Aachen,
 Sommerfeldstr. 16, 52074 Aachen, Germany
 E-mail: uttiya.sarkar@cern.ch

Identification of hadronic jets originating from heavy-flavor quarks is extremely important to several physics analyses in High Energy Physics, such as studies of the properties of the top quark and the Higgs boson, and searches for new physics. Recent algorithms used in the CMS experiment were developed using state-of-the-art machine-learning techniques to distinguish jets emerging from the decay of heavy flavour (charm and bottom) quarks from those arising from light flavour (up and down) quarks. Increasingly complex deep neural network architectures, such as graphs and transformers, have helped achieve unprecedented accuracies in jet tagging. New advances in tagging algorithms, along with new calibration methods using flavour-enriched selections of proton-proton collision events, allow us to estimate flavour tagging performances with the CMS detector during early Run 3 of the LHC.

42nd International Conference on High Energy Physics (ICHEP2024)
 16-24 July 2024
 Prague, Czech Republic

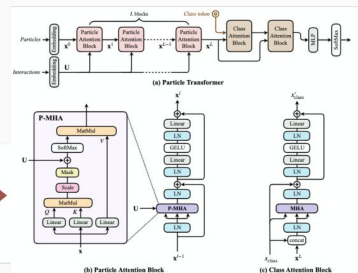
*Speaker

© Copyright owned by the author(s) under the terms of the Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International license. CC BY-NC-ND 4.0. <https://openaccess.cern.ch>

Transformers Inclusive Flavour Tagging

- Two separate IFT Transformer, B_s^0 and B^0 , trained on $B_s^0 \rightarrow D_s^- \pi^+$ and $B^0 \rightarrow J/\psi \phi$ simulations, and $B^0 \rightarrow J/\psi K^{*0}$
- Event collision represented as **point clouds**. Particle attention blocks learns contextual relations among tracks and captures both local and global event information.
- One prediction per event, no combination of decisions from SS or OS taggers [CMS-PAS-BPH-26-005]

$$U = \begin{aligned} \Delta &= \sqrt{(y_a - y_b)^2 + (\phi_a - \phi_b)^2}, \\ k_T &= \min(p_{T,a}, p_{T,b}) \Delta, \\ z &= \min(p_{T,a}, p_{T,b}) / (p_{T,a} + p_{T,b}), \\ m^2 &= (E_a + E_b)^2 - \|\mathbf{p}_a + \mathbf{p}_b\|^2. \end{aligned}$$



ParticleTransformer. Qu et al 2024. Particle Transformer for Jet Tagging.





VITA TOTA DISCENDA EST

Summary

Raymundus Lullus training program in advanced computing



UNIVERSIDADE DA CORUÑA



USC
UNIVERSIDADE DE SANTIAGO
DE COMPOSTELA



Universidad de Oviedo



IFCA



URL



Ciemat
Centro de Investigaciones
Energéticas, Medioambientales
y Tecnológicas



PIC

Artemisa



VITA TOTA DISCENDA EST

Summary

- The Transformer is the most important AI architecture of our time
- Its interface is next-token prediction, but it is not a probability machine
- It thinks and reasons. With the J-lens we can watch it happen
- To use its intelligence, what matters is how we talk to it
 - Vision is a matter of interface: Nano Banana
 - Knowledge is a matter of interface: RAG

Raymundus Lullus training program in advanced computing





VITA TOTA DISCENDA EST

Thanks!

Raymundus Lullus training program in advanced computing





VITA TOTA DISCENDA EST

Hands-on session

Raymundus Lullus training program in advanced computing



UNIVERSIDADE DA CORUÑA



USC
UNIVERSIDADE DE SANTIAGO
DE COMPOSTELA



Universidad de Oviedo



IFCA



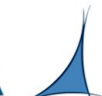
URL



Ciemat
Centro de Investigaciones
Energéticas, Medioambientales
y Tecnológicas



GAPA



Artemisa



VITA TOTA DISCENDA EST

Hands-on session

45' hands-on

j-lens experiments

- Demonstrate the intelligence
- Try it yourself

RAG exercise

- Demonstrate how to give Transformer new knowledge

Raymundus Lullus training program in advanced computing



UNIVERSIDADE DA CORUÑA



UNIVERSIDADE DE SANTIAGO DE COMPOSTELA



Universidad de Oviedo



URL



Centro de Investigaciones Energéticas, Medioambientales y Tecnológicas



IPARCOS

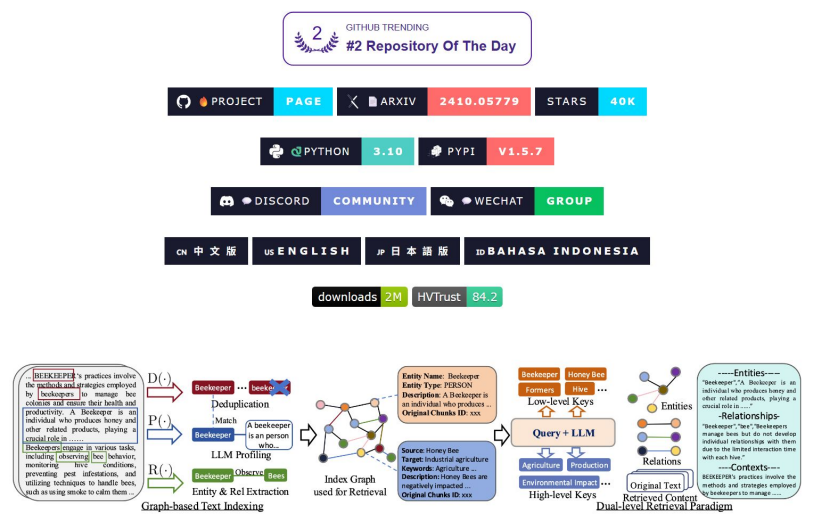


Artemisa



RAG exercise

LightRAG: Simple and Fast Retrieval-Augmented Generation



github.com/hkuds/lightrag

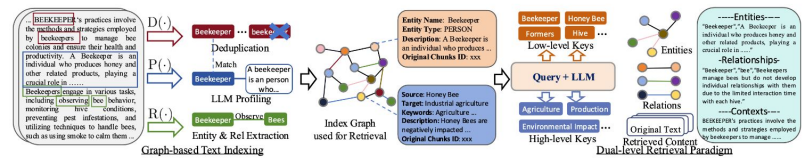


Figure 1: Overall architecture of the proposed LightRAG framework.

Raymundus Lullus training program in advanced computing



VITA TOTA DISCENDA EST

Hands-on session

All the code used today is on



`github.com/JiahuiZhuo/jlens-explorer`

`github.com/JiahuiZhuo/lightrag-pdf-kb`

Raymundus Lullus training program in advanced computing

