

Multimodal Foundation Models in the VLAM Program



Tanguy Coenen
Scientific Director at imec

Let's talk about...

- The VLAM program
- Multimodal AI in media
 - Use case: agentic AI in media
 - Use case: virtual production
- Challenges



The VLAM program



The Digital Sovereignty Imperative

Need for Attuned AI

Critical need for culturally and linguistically attuned AI in a global landscape dominated by large international models.

Risk of Bias

Risk of cultural erosion and linguistic bias if we rely solely on dominant, non-local models.



VLAM

Develop robust, culturally attuned foundation models to serve as the backbone for the Flemish ecosystem.



The Foundation: Rich Media Archives

Data is the Core

Unique value from rich digital archives (text, audio, video) of Flemish media houses.

This data enriches open-source international models.

Goal: Create AI that reflects the region's linguistic and cultural identity.



Strategy: "Adapt, Don't Build" => work on post-training and inference

Building from Scratch

- Extremely high cost (billions)
- Requires massive compute resources
- Slow time-to-market
- Difficult to keep pace with global leaders

The VLAM Way: Adapt

- Cost-efficient and fast
- Uses existing State-of-the-Art (SOTA) models as a base
- Focuses resources on fine-tuning with unique Flemish data
- Ensures technological relevance



Core Component: The Foundation Base Model



Base Models: We start with powerful, open-source international foundation models as our technological core.



Flemish Data: We enrich this base with our unique, high-quality multimodal media archives.



Process: A continuous process of fine-tuning, retrieval-augmentation, and grounding ensures deep cultural relevance and accuracy.



Multimodal AI in media



Input modality

Text

This is a cat

Code

```
{  
  "type": "cat",  
  "Description": "cute"  
}
```

Image



Audio



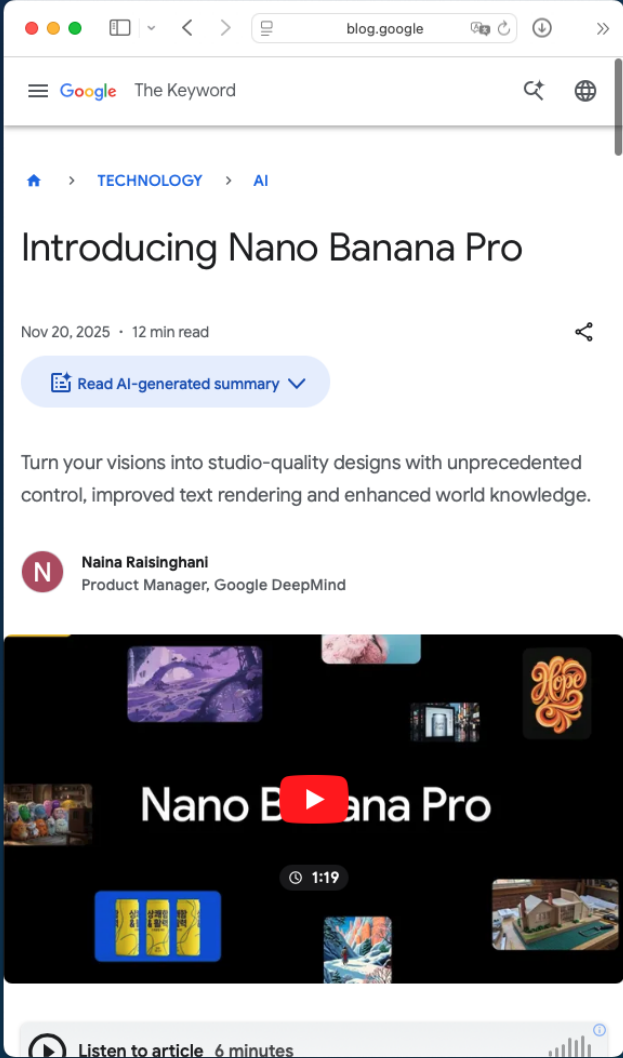
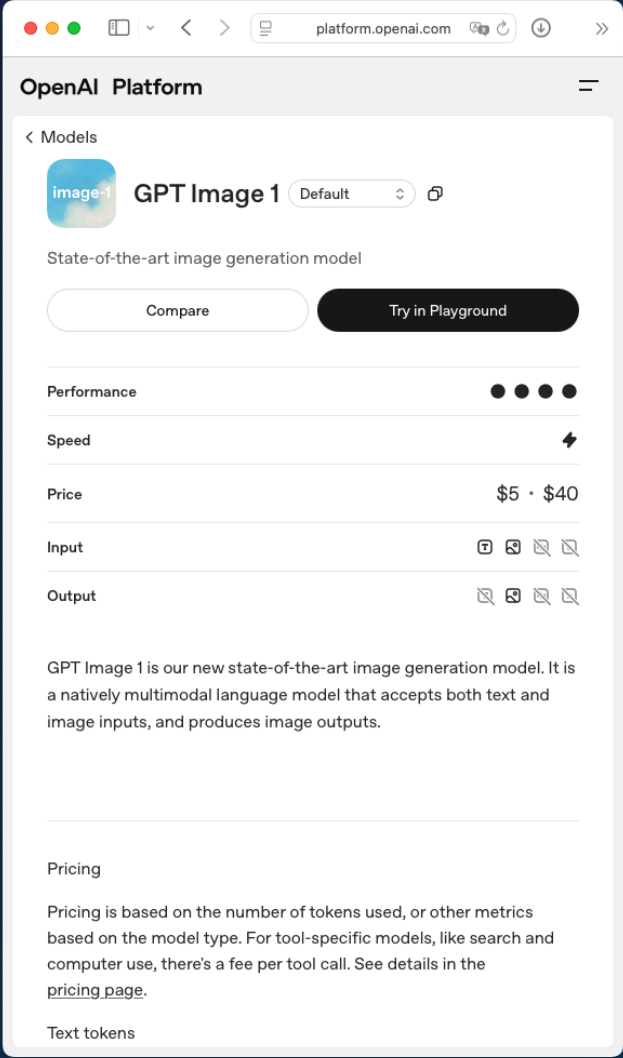
Multimodal
model

Output modality

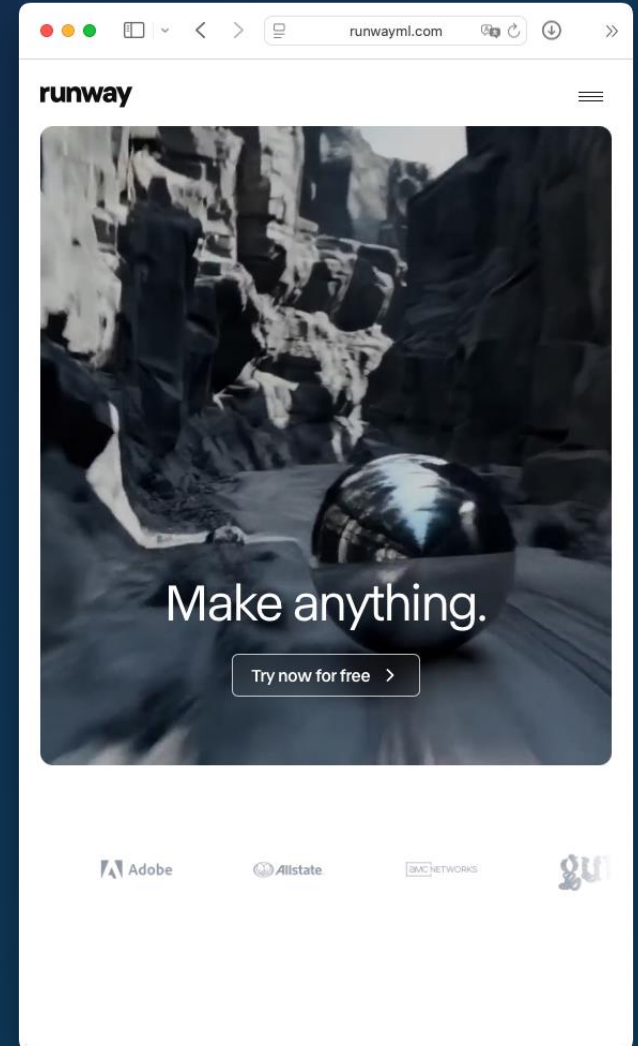
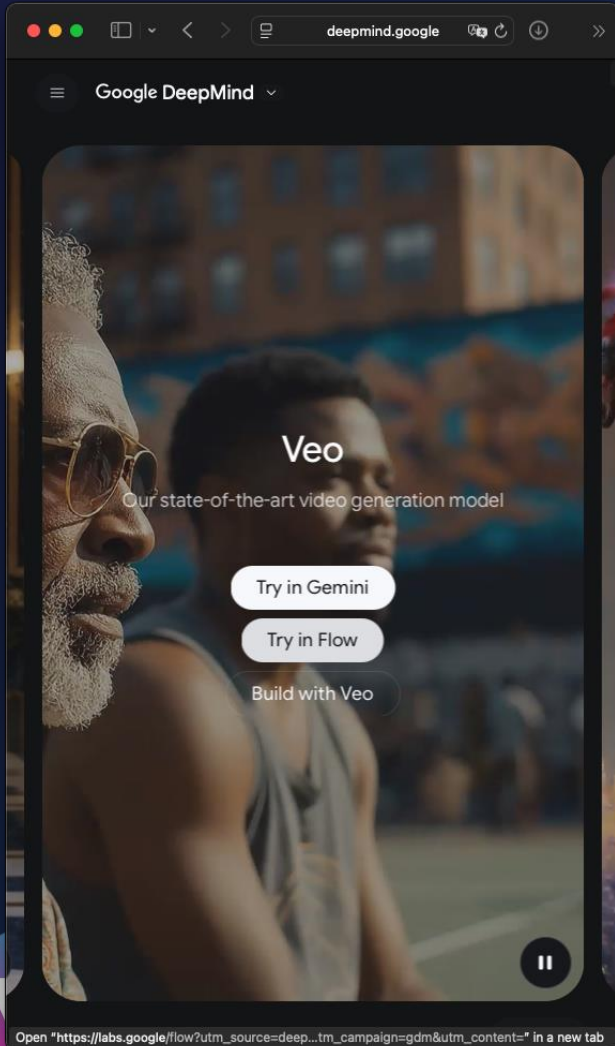
It's pixelated!

```
{  
  "status": "accepted"  
}
```

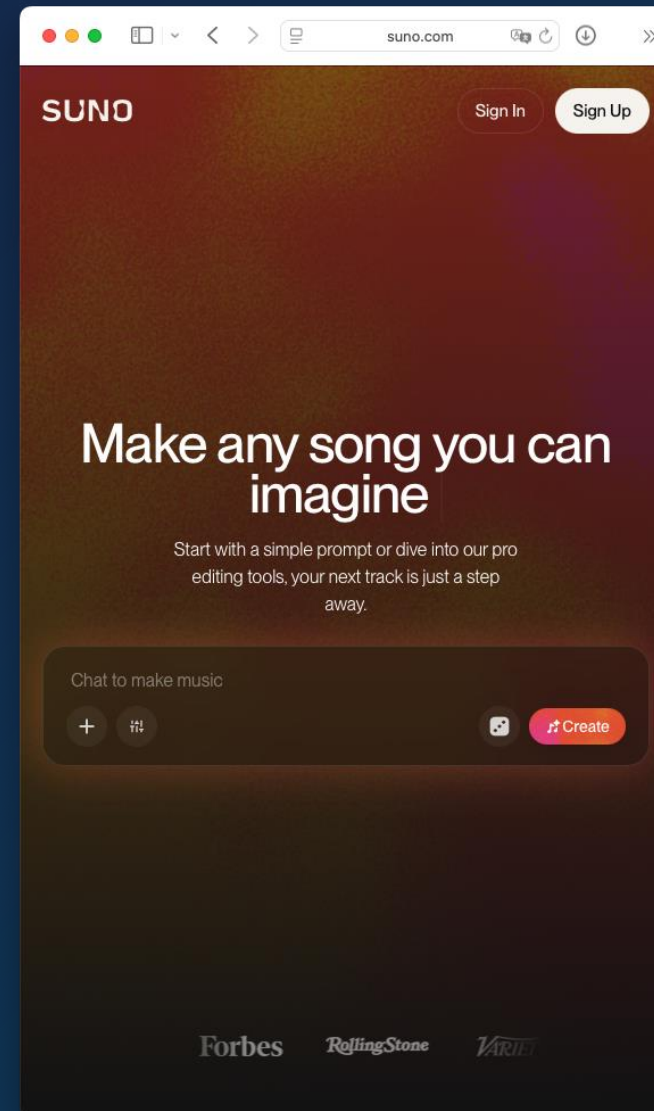
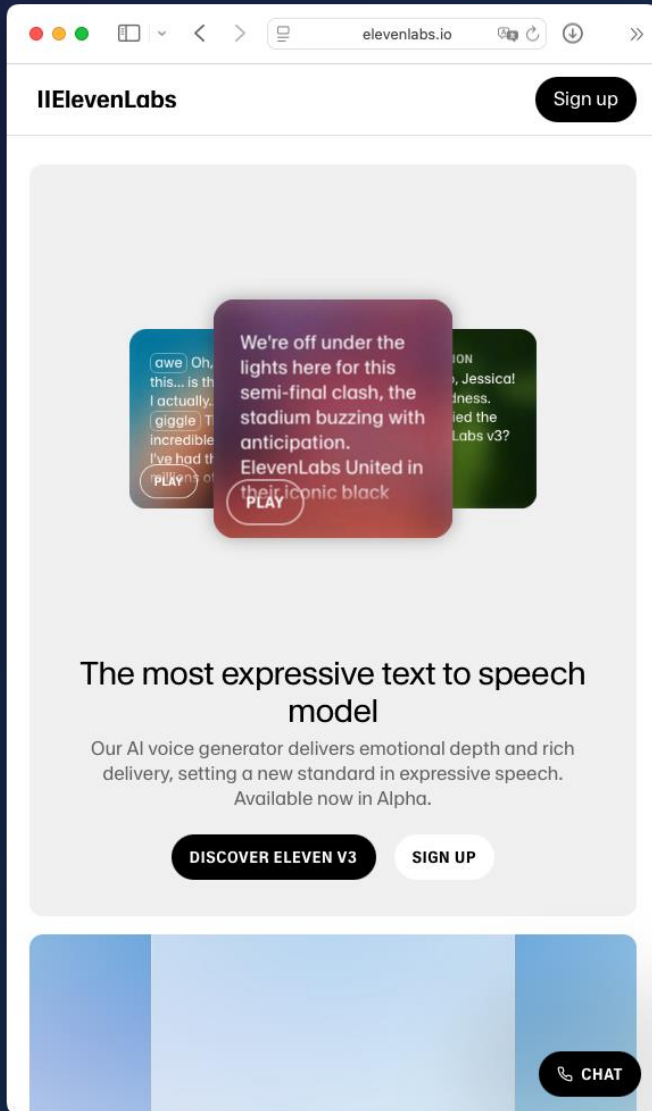
Image



Video



Audio





The screenshot shows the Hugging Face interface for the model **baidu/ERNIE-4.5-VL-28B-A3B-Thinking**. The header includes the Hugging Face logo, a search bar, and navigation links like "Model card", "Files", "xet", and "Community". Below the header, there are tabs for different interfaces: "Image-Text-to-Text", "Transformers", "Safetensors", "English", "Chinese", and "ernie4_5_moe_vl". There are also filters for "feature-extraction", "ERNIE4.5", "conversational", "custom_code", and "License: apache-2.0". A "Deploy" button and a "Use this model" button are visible.

The main content area features a section titled "Introducing ERNIE-4.5-VL-28B-A3B-Thinking: A Breakthrough in Multimodal AI" with a "Demo" link. This section highlights the model's capabilities in multimodal reasoning and its training process. To the right of this text is a "Downloads last month" chart showing a peak at 38,991 downloads.

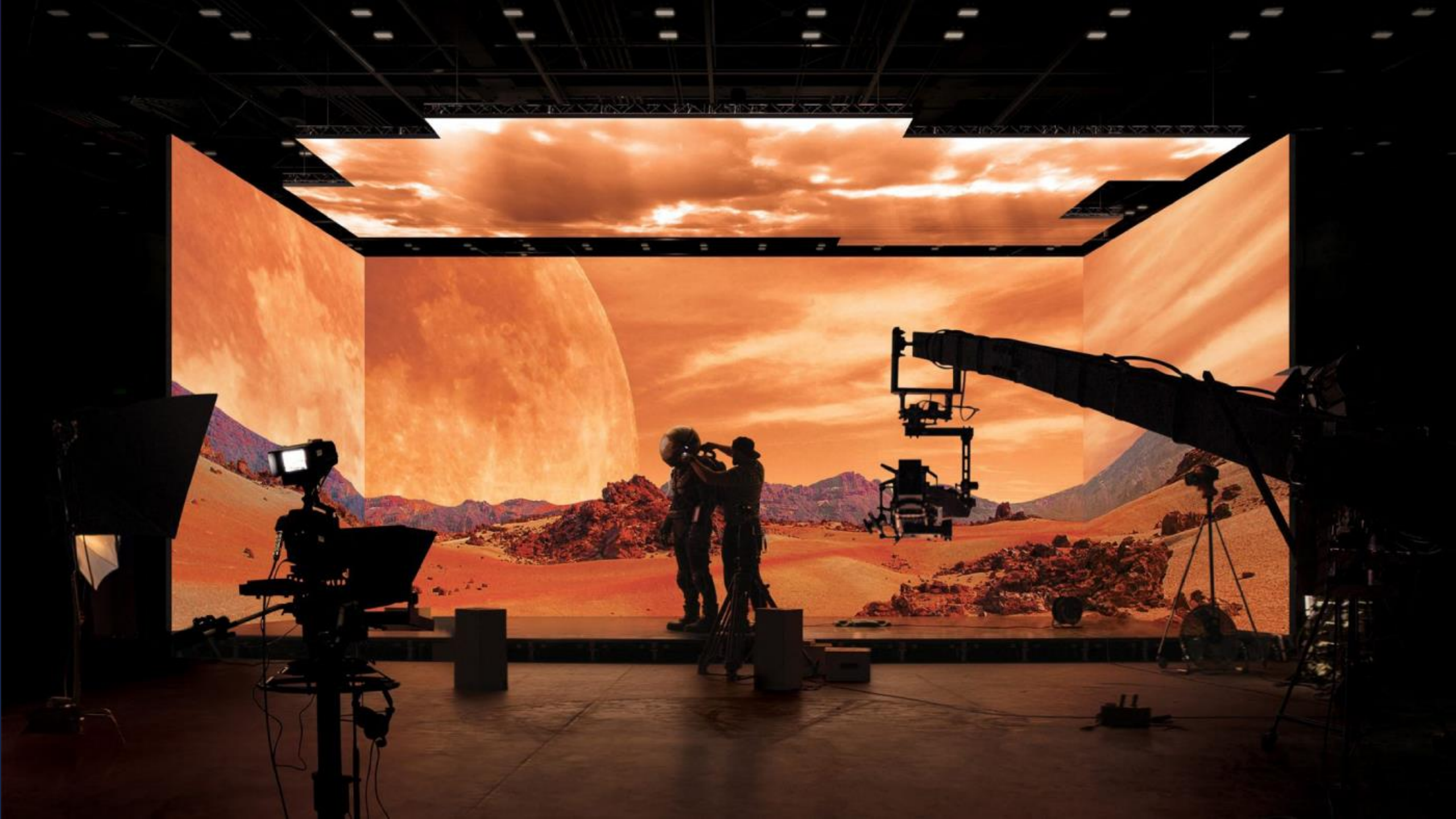
Below the introduction, there are sections for "Model Highlights", "Inference Providers", and "Model tree for baidu/ERNIE-4.5-V...". The "Model Highlights" section describes the model's architecture and training. The "Inference Providers" section lists providers like "Image-Text-to-Text" and notes that the model isn't deployed by any provider. The "Model tree" section shows quantizations and spaces using the model.

At the bottom, there are social media links for Chat, GitHub, Blog, Discord, PaddlePaddle, and ErnieforDevs, along with the license "Apache2.0".

Use case 1: Agentic AI in media



Use case 2: Virtual production





Challenges



Some of the multimodal challenges we plan to work on in VLAM

- Authenticity tracking
- Data value attribution
- Multimodal RAG
- Narrative structure
- Emotional connection
- Collaboration between The Netherlands and Flanders



Thank You

tanguy.coenen@imec.be

