

From Volumes to Clinical Decision

Five conditions to fill the gap between medical 3D vision and the clinic

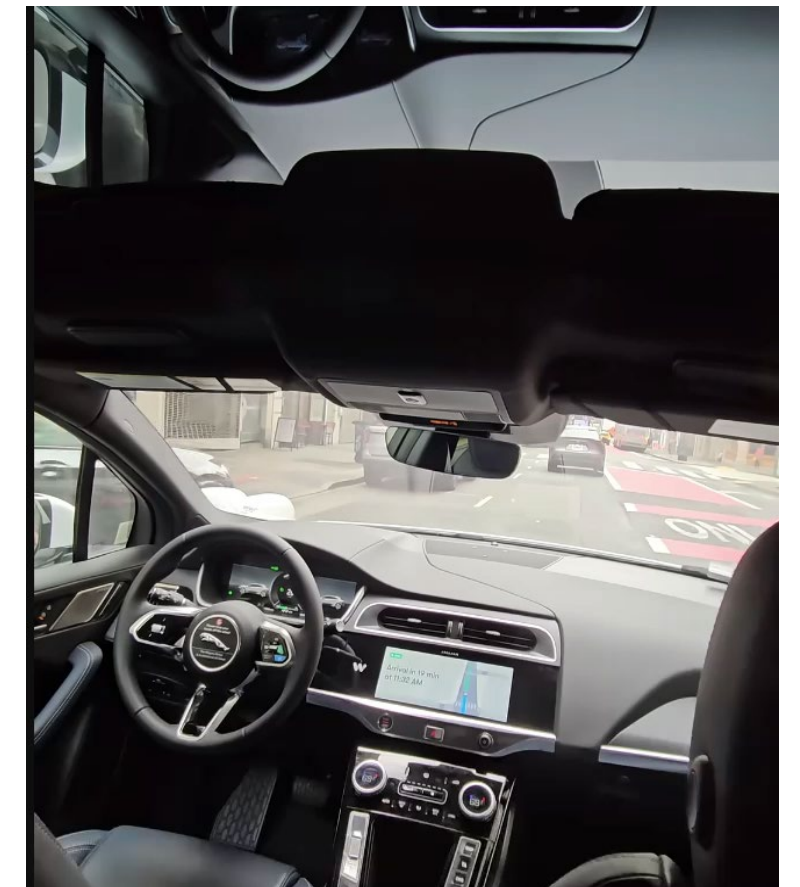
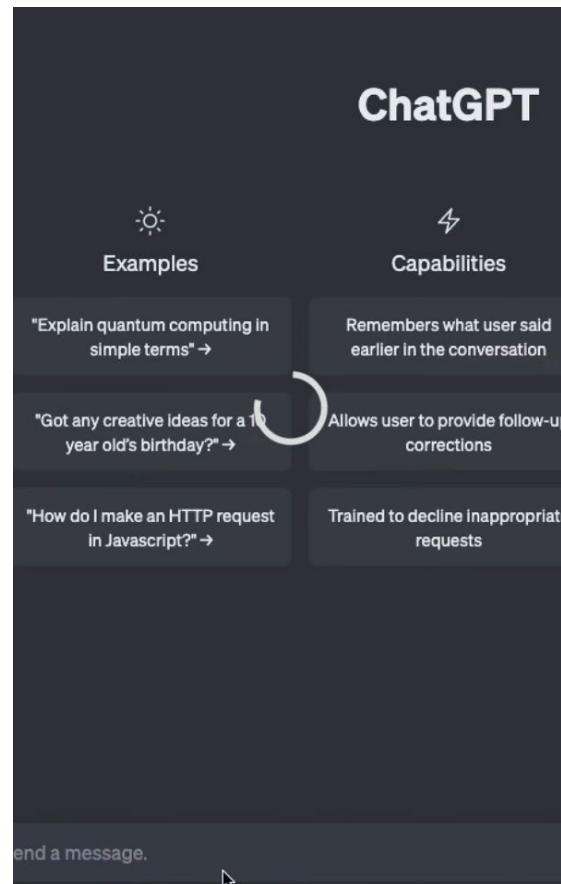
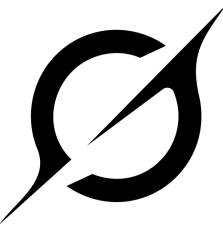
Shared Data · Reliable Supervision · Robust Models
Actionable Outputs · Trusted Evaluation

Federico Bolelli

AlmageLab — University of Modena and Reggio Emilia · federico.bolelli@unimore.it



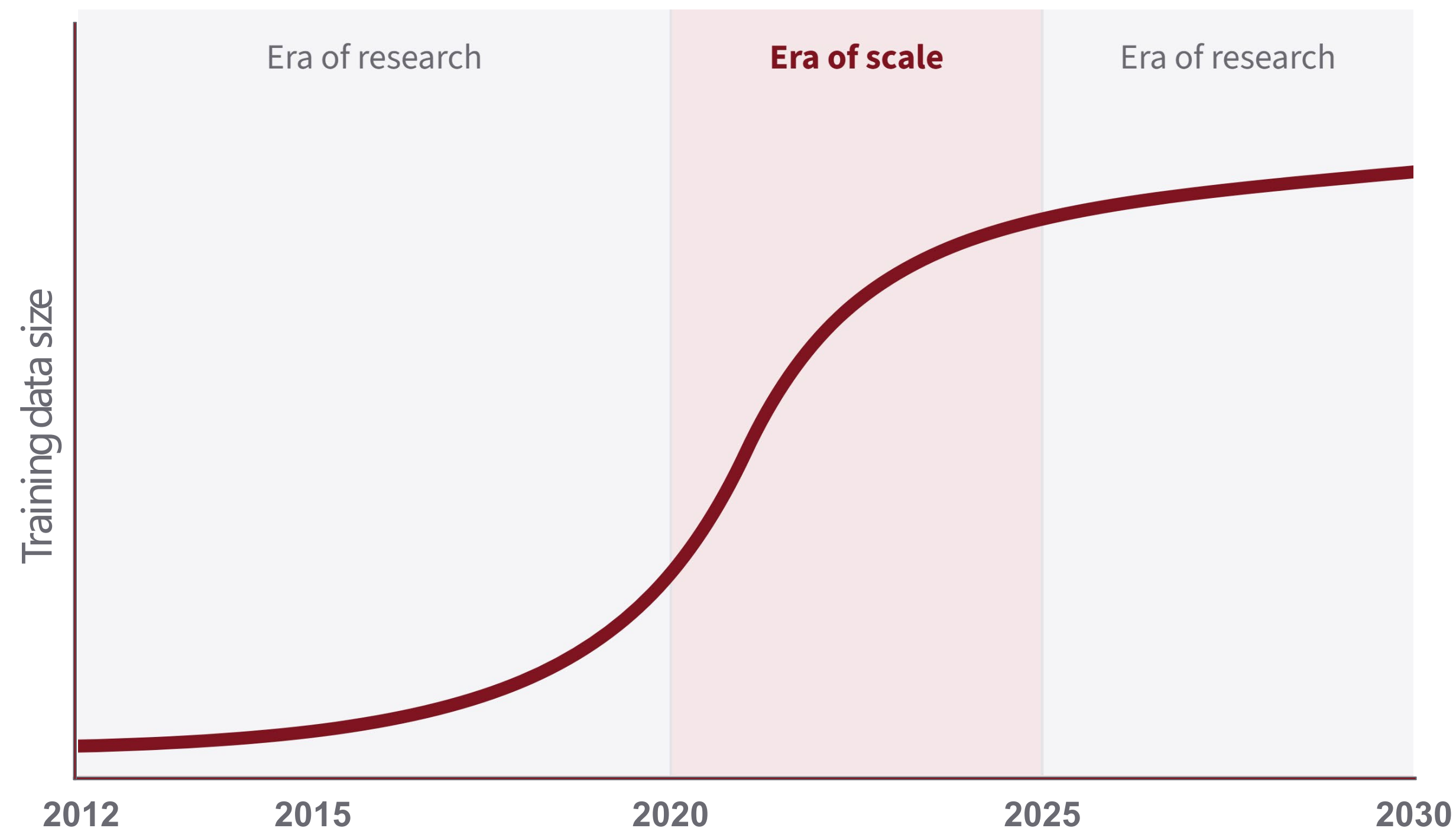
UNIMORE
UNIVERSITÀ DEGLI STUDI DI
MODENA E REGGIO EMILIA



Video not available in the PDF!

Why Medical 3D Vision Has No ImageNet Moment (Yet)?

General models crossed into the **era of scale**.



VISION-LANGUAGE **650×**

SigLIP 2 10B

BiomedCLIP 15.3M

SEGMENTATION **700×**

SAM 1.1B

MedSAM 1.57M

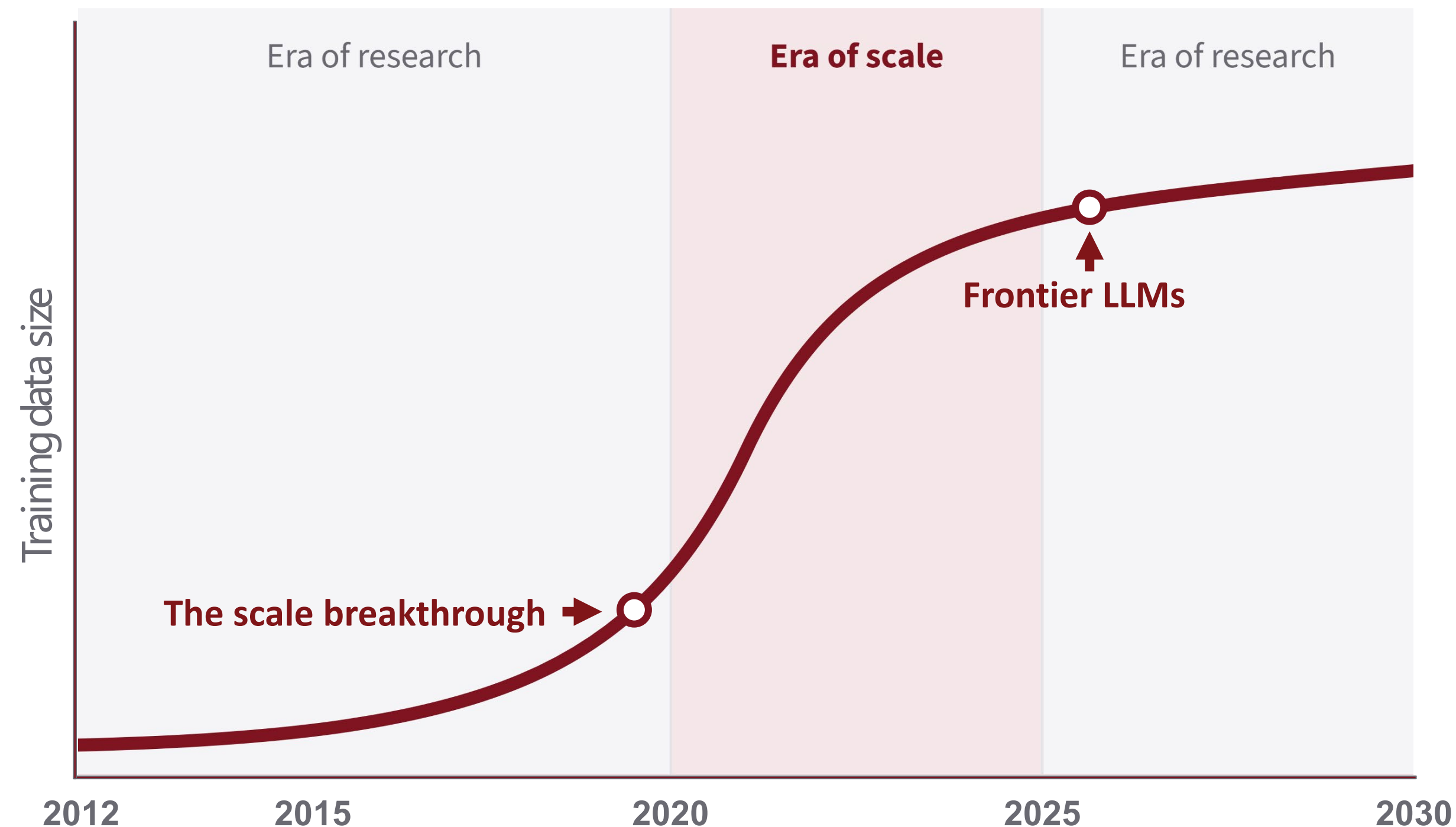
CLASSIFICATION **11×**

ImageNet-21k 14.2M

RadImageNet 1.35M

Why Medical 3D Vision Has No ImageNet Moment (Yet)?

General models crossed into the **era of scale**.



VISION-LANGUAGE **650×**

SigLIP 2 10B

BiomedCLIP 15.3M

SEGMENTATION **700×**

SAM 1.1B

MedSAM 1.57M

CLASSIFICATION **11×**

ImageNet-21k 14.2M

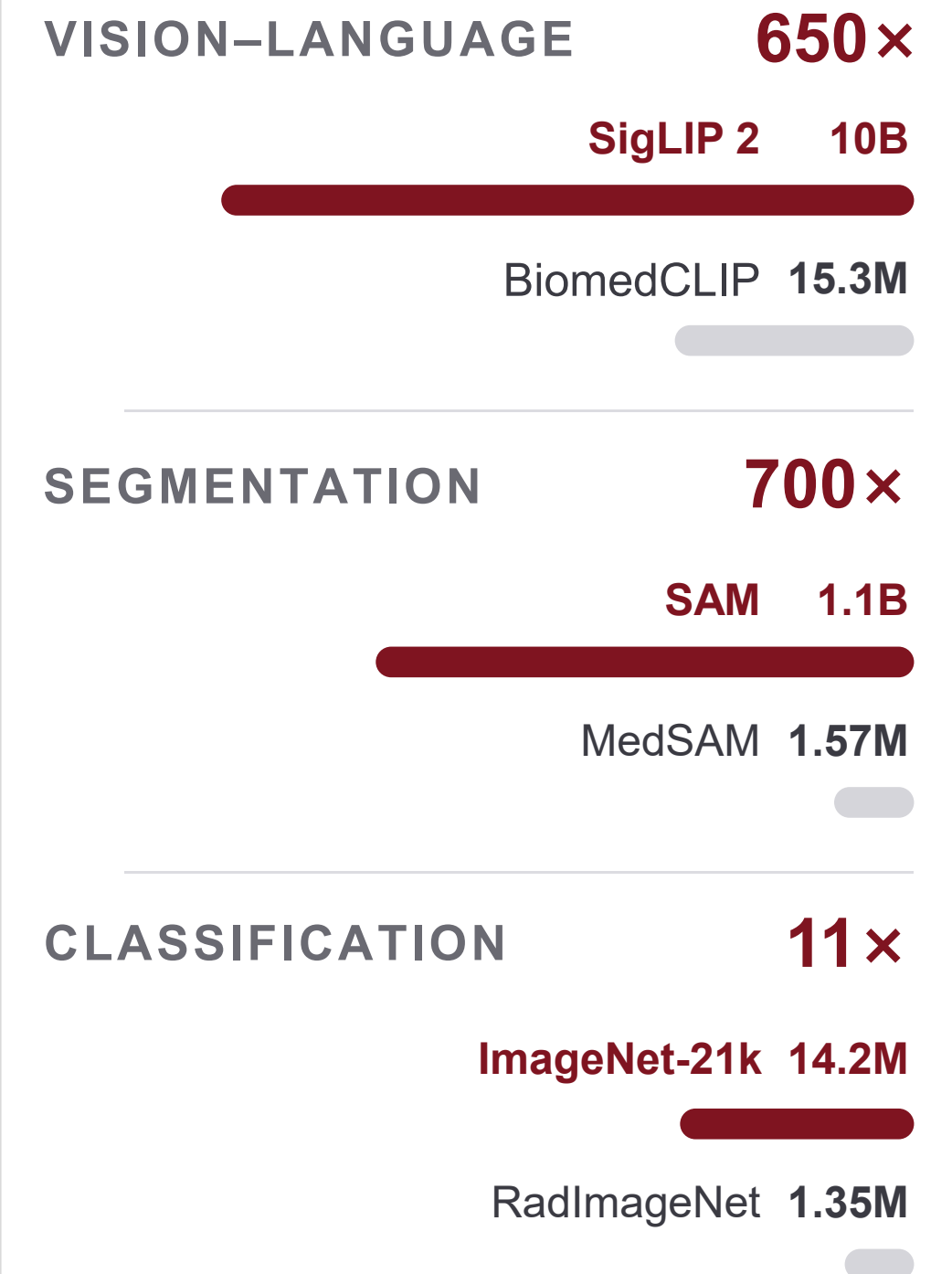
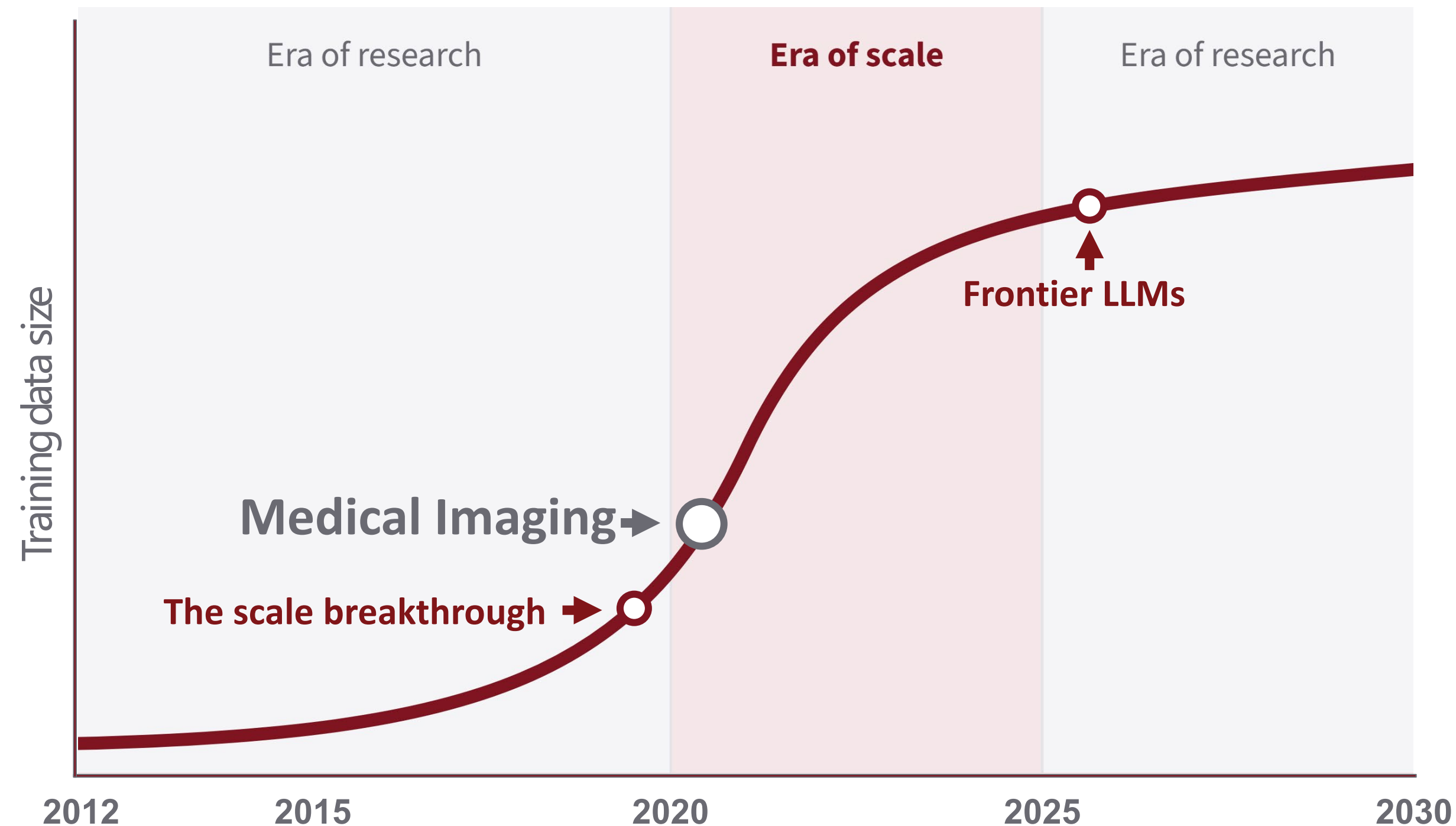
RadImageNet 1.35M

Counts as reported by the authors: WebLI (SigLIP 2), PMC-15M (BiomedCLIP), SA-1B (SAM), MedSAM, ImageNet-21k, RadImageNet.

[1] Curtis P. Langlotz - Toward a Radiology Foundation Model

Why Medical 3D Vision Has No ImageNet Moment (Yet)?

General models crossed into the **era of scale**.



Counts as reported by the authors: WebLI (SigLIP 2), PMC-15M (BiomedCLIP), SA-1B (SAM), MedSAM, ImageNet-21k, RadImageNet.

[1] Curtis P. Langlotz - Toward a Radiology Foundation Model.

WHERE IT ALREADY WORKS

PATHOLOGY — BREAST, WSI

0.994
algorithm

0.884
pathologist

AUC of the best algorithm against the best pathologist in a time-constrained diagnostic setting. Seven of 32 algorithms exceeded a panel of 11 pathologists.

DERMOSCOPY — MELANOMA

0.86
CNN

0.79
dermatologists

ROC AUC of a CNN against the mean of 58 dermatologists in a reader study.

BREAST HISTOLOGY

99%
panel of four

85%
single observer

Accuracy on benign versus malignant discrimination on previously unseen slides; pooling four independent observers raises it to 99%.

[1] Ehteshami Bejnordi, Babak, et al. "Diagnostic Assessment of Deep Learning Algorithms for Detection of Lymph Node Metastases in Women With Breast Cancer" *Jama* 2017.

[2] Haenssle, Holger A., et al. "Man against machine: diagnostic performance of a deep learning convolutional neural network for dermoscopic melanoma recognition in comparison to 58 dermatologists." *Annals of Oncology* 2018.

[3] Levenson et al., "Pigeons (*Columba livia*) as Trainable Observers of Pathology and Radiology Breast Cancer Images" *PLoS ONE* 2015.

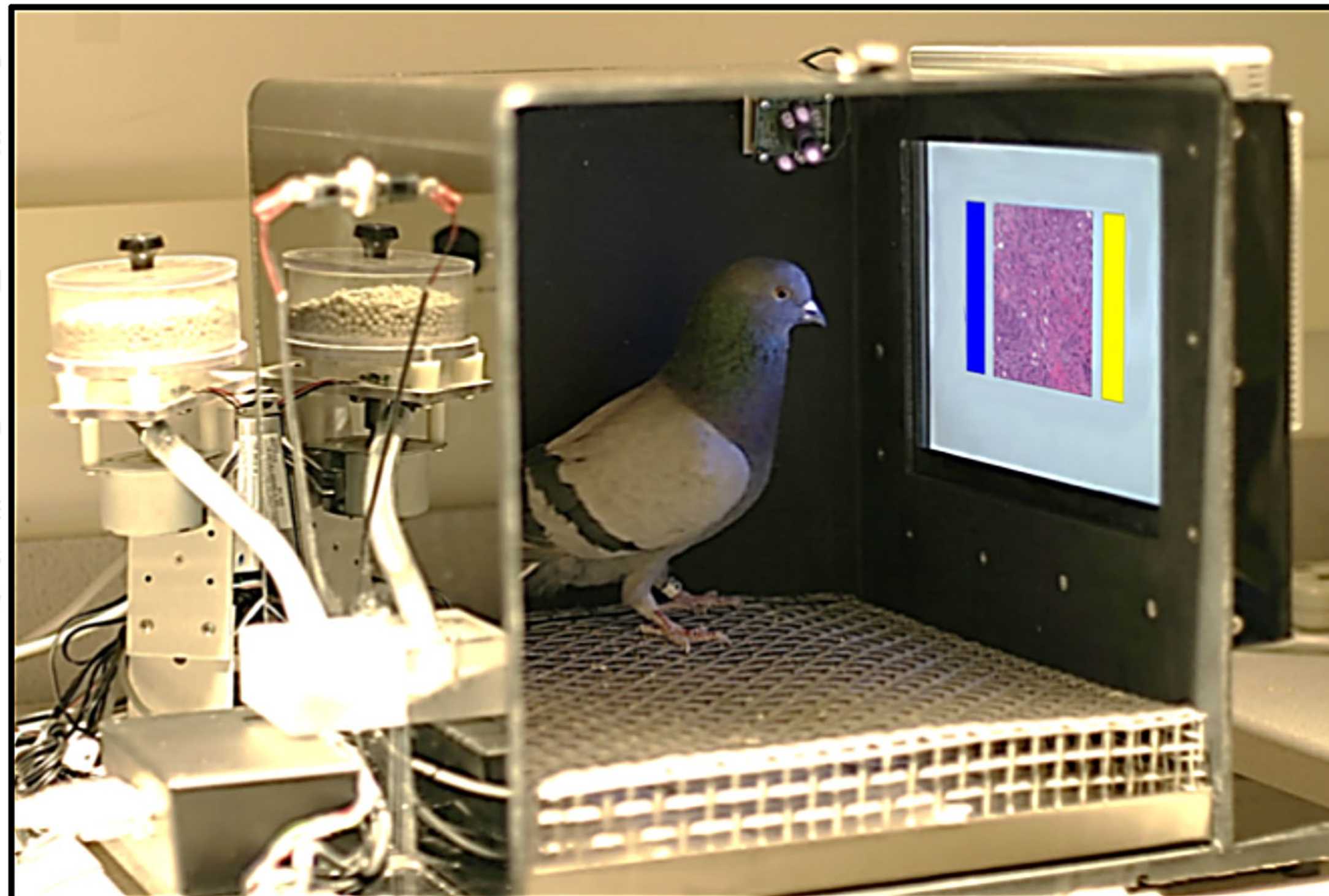
WHERE IT ALREADY WAS

PATHOLOGY — BREAST

0.994

0.884

AUC of the best algorithm a pathologist in a time-constrained setting. Seven of 32 algorithms panel of 11 pathologists.



PATHOLOGY

panel of four

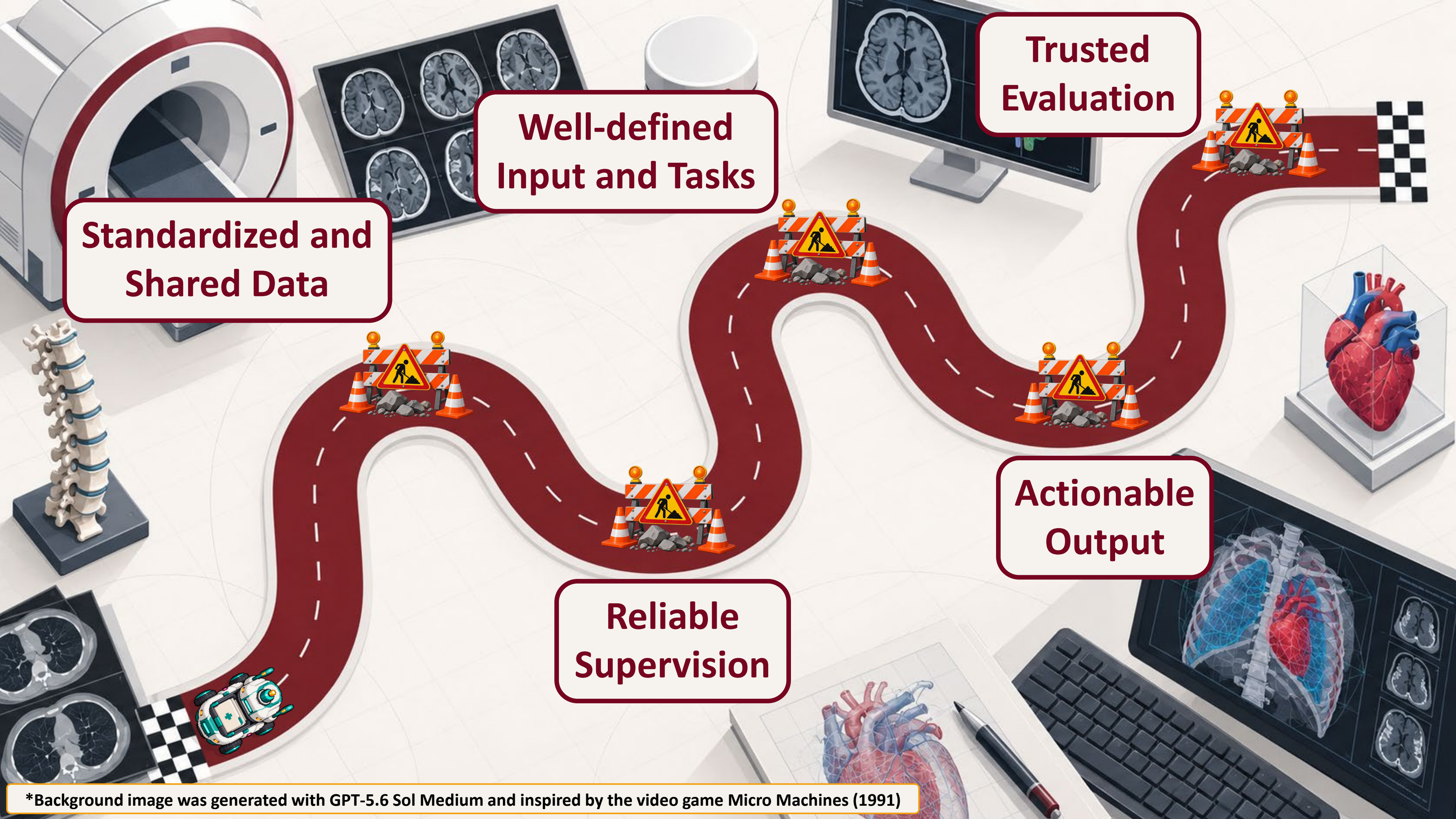
single observer

benign versus malignant
previously unseen slides;
independent observers raises it

[1] Ehteshami Bejnordi, Babak, et al. "Diagnostic Assessment of Deep Learning Algorithms for Detection of Lymph Node Metastases in Women With Breast Cancer" *Jama* 2017.

[2] Haenssle, Holger A., et al. "Man against machine: diagnostic performance of a deep learning convolutional neural network for dermoscopic melanoma recognition in comparison to 58 dermatologists." *Annals of Oncology* 2018.

[3] Levenson et al., "Pigeons (*Columba livia*) as Trainable Observers of Pathology and Radiology Breast Cancer Images" *PLoS ONE* 2015.



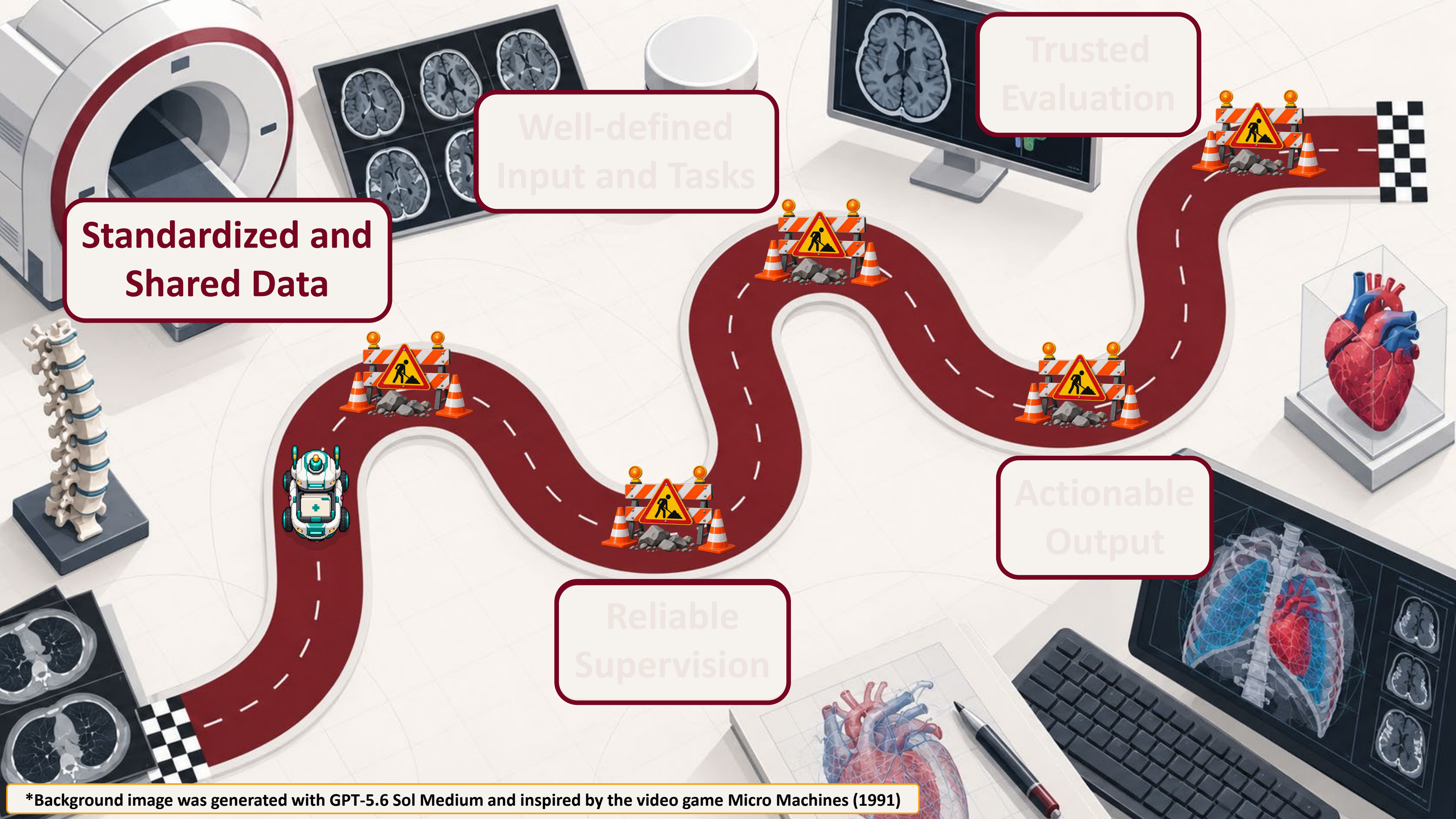
**Standardized and
Shared Data**

**Well-defined
Input and Tasks**

**Trusted
Evaluation**

**Actionable
Output**

**Reliable
Supervision**



**Standardized and
Shared Data**

**Well-defined
Input and Tasks**

**Trusted
Evaluation**

**Actionable
Output**

**Reliable
Supervision**

MEDICAL IMAGING BENCHMARKS

Three layers now make methods more comparable.

01

Shared task

Same cohort, image inputs, reference labels and train/test split.

BraTS · CAMELYON · Decathlon

02

Controlled evaluation

Hidden test labels, central scoring and a fixed submission protocol.

MICCAI challenge ecosystem

03

Standardized evidence

Task-aware metric selection, ranking rules and transparent reporting.

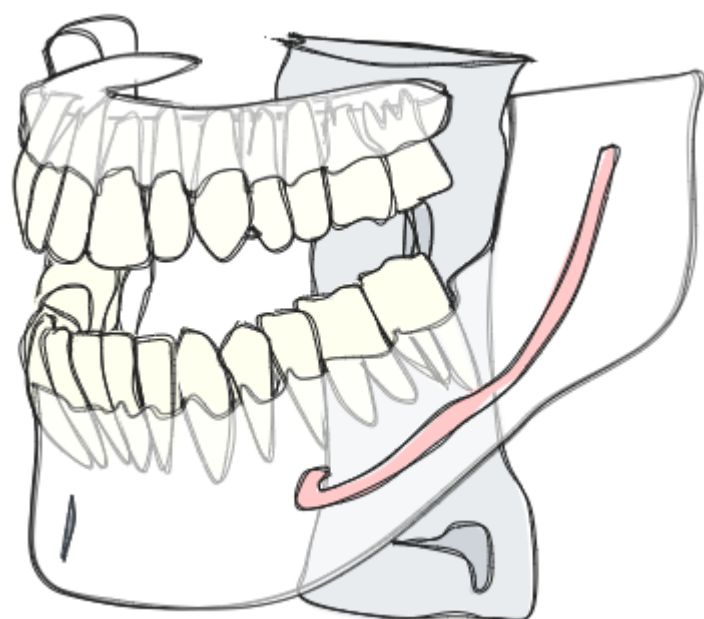
BIAS · Metrics Reloaded

[1] L. Maier-Hein, et al. “Why rankings of biomedical image analysis competitions should be interpreted with care.” *Nature Communications* 2018.

[2] L. Maier-Hein, et al. “Metrics reloaded: recommendations for image analysis validation.” *Nature Methods* 2024.



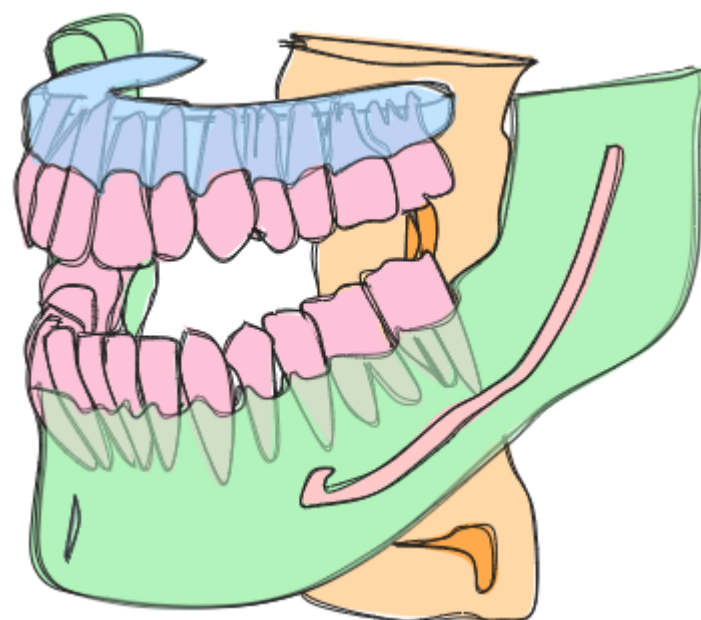
1 class



2023



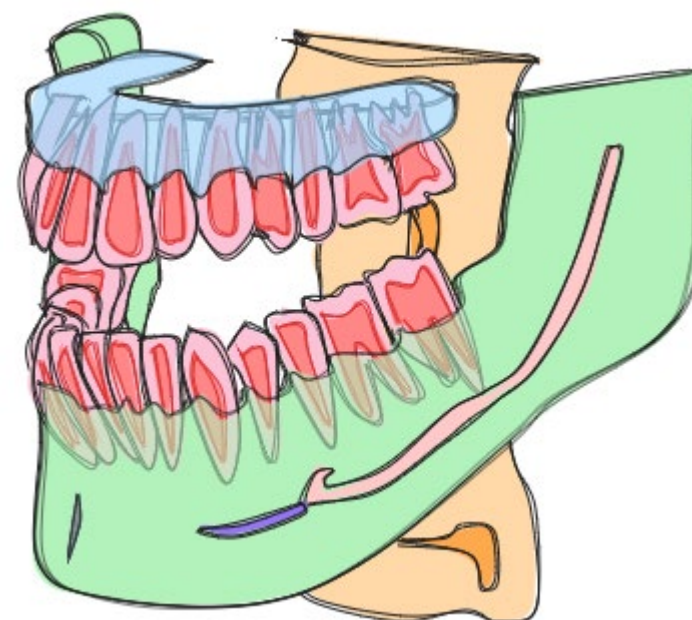
42 classes



2024



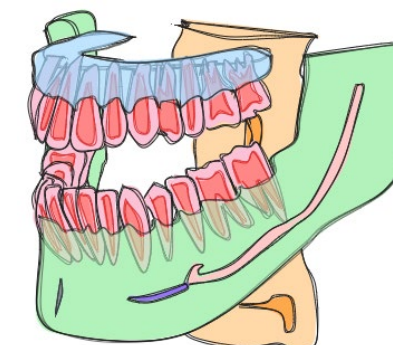
77 classes



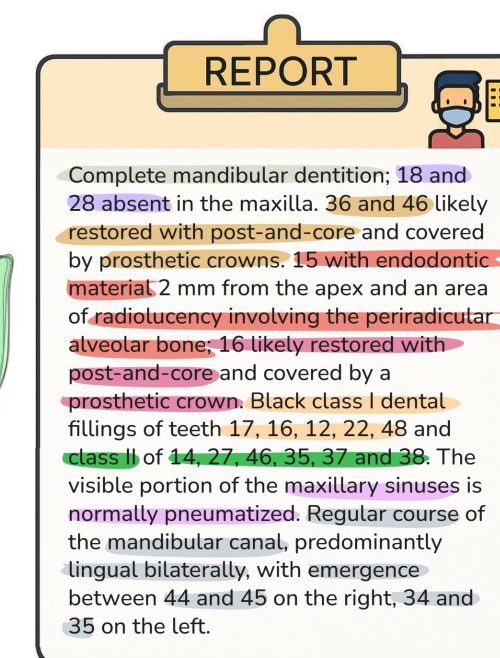
2025



77 classes + Reports



2026



[1] Bolelli F., et al. "Segmenting the Inferior Alveolar Canal in CBCT Volumes: The ToothFairy Challenge" IEEE Transactions on Medical Imaging 2024.

[2] Bolelli F., et al. "Multi-structure segmentation in CBCT volumes: The ToothFairy2 challenge." Medical Image Analysis 2026.

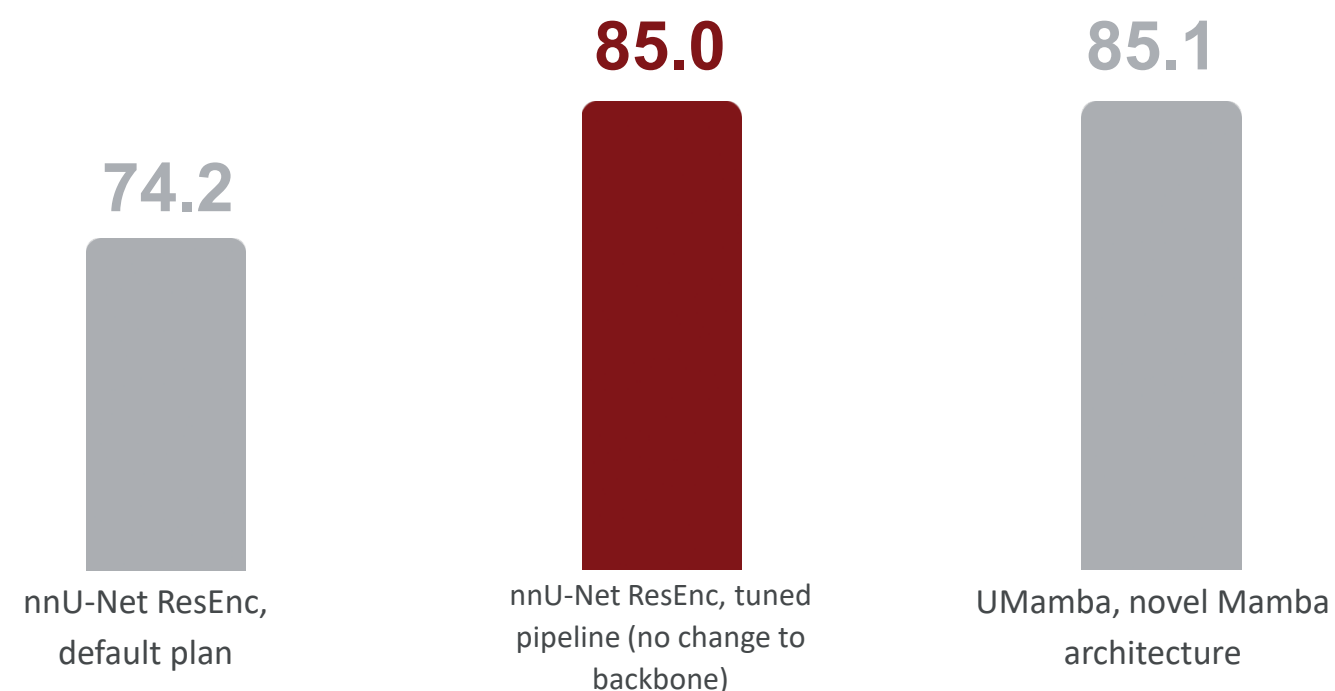
[3] Lumetti L., et al. "ToothFairy3: Scaling CBCT Maxillofacial Segmentation to 77 Classes with U-Mamba2." MICCAI 2026.

[4] Lumetti L., et al. "Ontology-Grounded Structured Prediction for Dental CBCT Reporting." MICCAI 2026.

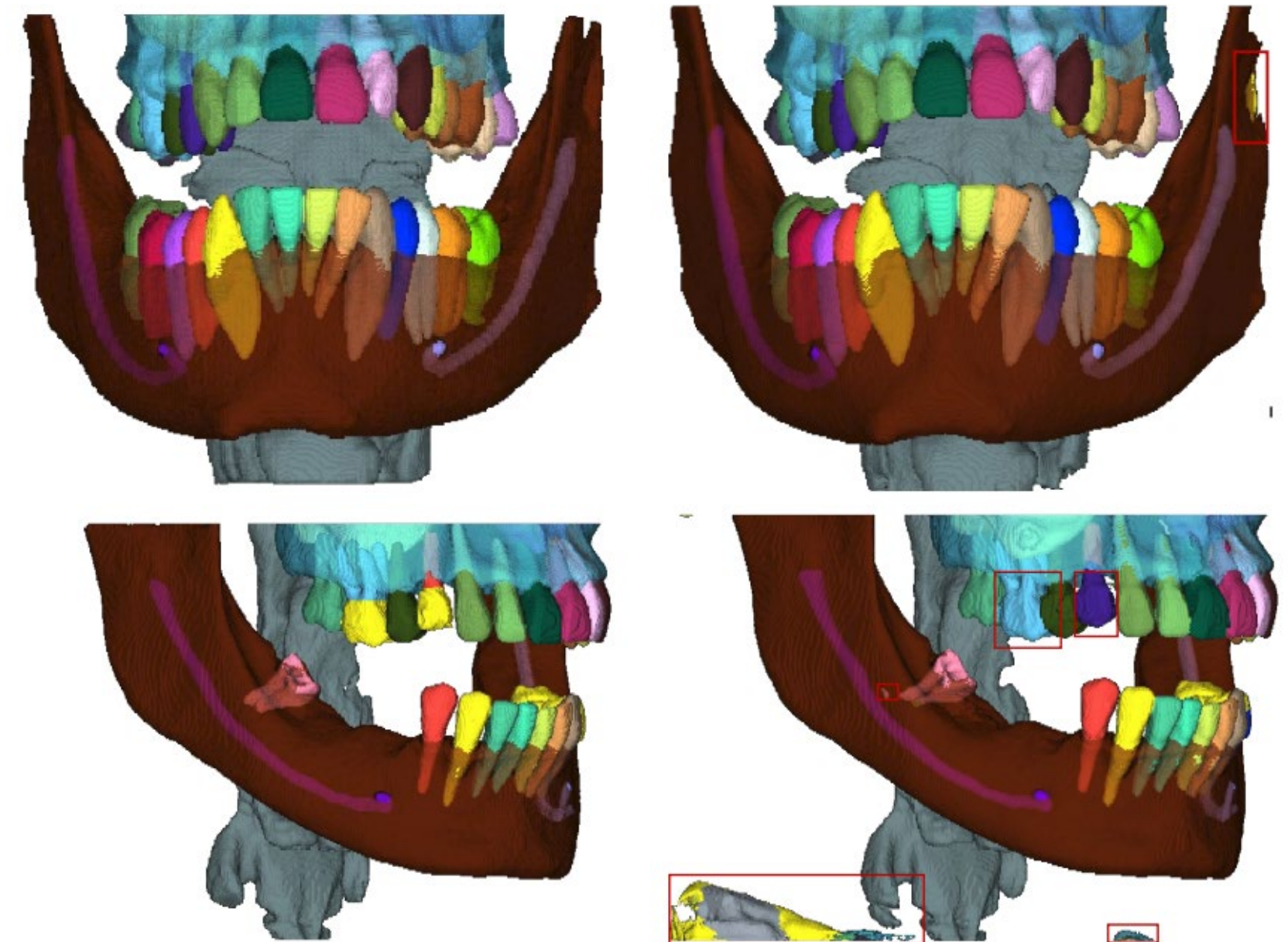
When methods are evaluated under the same protocol, strong nnU-Net-based pipelines remain extremely difficult to beat.

- This does not mean architectures do not matter.
- Once the benchmark is controlled, preprocessing, post-processing, training strategy and data quality can become at least as important as changing the backbone.

TOOTHFAIRY2 BENCHMARK — DICE SCORE (DSC, %)



QUALITATIVE RESULT



(a) Ground-Truth

(b) nnU-Net ResEnc

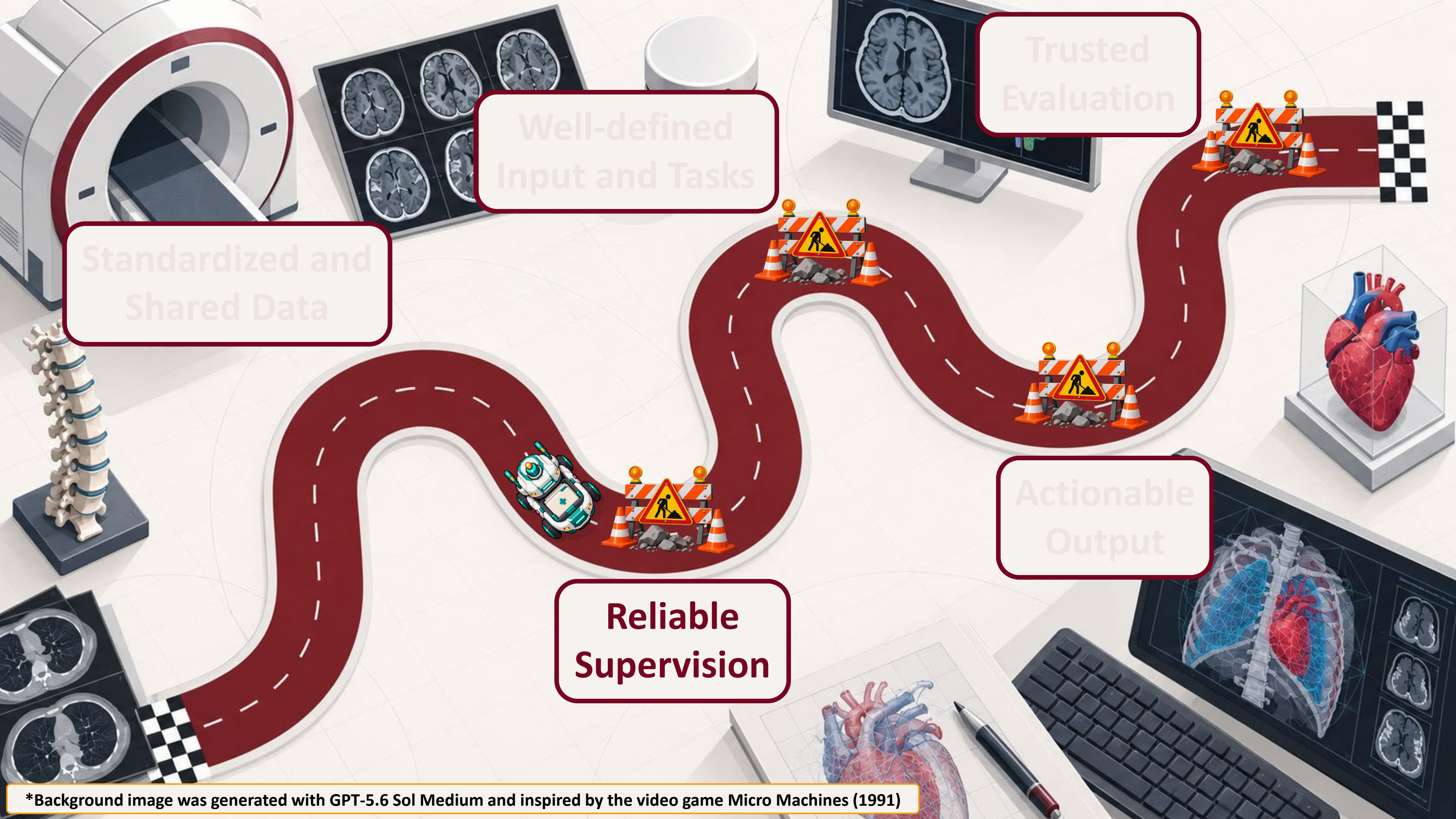
[1] Bolelli, F., Marchesini K., et al. "Segmenting Maxillofacial Structures in CBCT Volumes". CVPR 2025.

[2] Isensee, F., Wald, T., Ulrich, C., et al. nnU-Net revisited: A Call for Rigorous Validation in 3D Medical Image Segmentation. MICCAI 2024.



 ToothFairy2 CBCT dataset of 480 volumes with fine-grained annotations spanning 42 maxillofacial classes (teeth, ja... Segmentation MICCAI 2024 Jul 2024 1429	 ToothFairy3 CBCT dataset of 532 NIFTI volumes with comprehensive dental annotations covering 77 classes... Segmentation MICCAI 2025 MICCAI 2026 Jun 2025 1205	 ToothFairy CBCT dataset for inferior alveolar nerve and dental segmentation, released for the MICCAI 2023... Segmentation MICCAI 2023 Jun 2026 814	 ToothFairy4 Multimodal dataset of 625 CBCT cases paired with clinical reports in Italian and English, from multiple... Report generation MICCAI 2026 Jul 2026 479	 BarBeR Barcode benchmarking dataset of 8,748 real images with 9,818 polygon-annotated barcodes across... Detection Localization ICPR 2024 Jun 2026 406	 Bits2Bites Dataset of 200 paired upper/lower intra-oral 3D scans (STL) of dental arches collected by Italian... Classification ODIN 2025 Jul 2026 299	 Pulpy3D Dental CT dataset of 423 volumes with semantic and instance annotations across 19 classes... Segmentation MICCAI 2024 Dec 2024 280
 Bite2Text Orthodontic dataset of 2,000 patient cases pairing 3D intra-oral scans and 2D photographs with clinician... Report generation ECCV 2026 MICCAI 2026 Jun 2026 219	 Maxillo Maxillofacial CBCT dataset for inferior alveolar canal (IAC) segmentation, with 91 densely an... Segmentation CVPR 2022 IEEE Access 2022 Jun 2026 193	 ResiDual Industrial inspection dataset of 13,870 images grouped into 730 sequences of pharmaceutical vials... Classification Detection ICPR 2024 Jun 2026 48	 TesticulUS Collection of about 9,300 synthetic B-mode testicular ultrasound images (256x256 PNG) generated with... Classification Segmentation BMVC 2026 ICIAP 2025 Jun 2026 27	 IOP-Compass Dataset of 5,000 standardized intraoral photographs from 1,000 patients — five orthodontic views... Classification Segmentation ODIN 2026 Aug 2026 14	 ReportX A Clinician-Curated Brain Tumor MRI Report Dataset Report generation MICCAI 2026 Jun 2026 12	 CALHippo The Cellular Annotation Library for the Hippocampus Classification Reconstruction Segmentation MICCAI 2026 Jul 2026 5

Video not available in the PDF!



Well-defined
Input and Tasks

Trusted
Evaluation

Standardized and
Shared Data

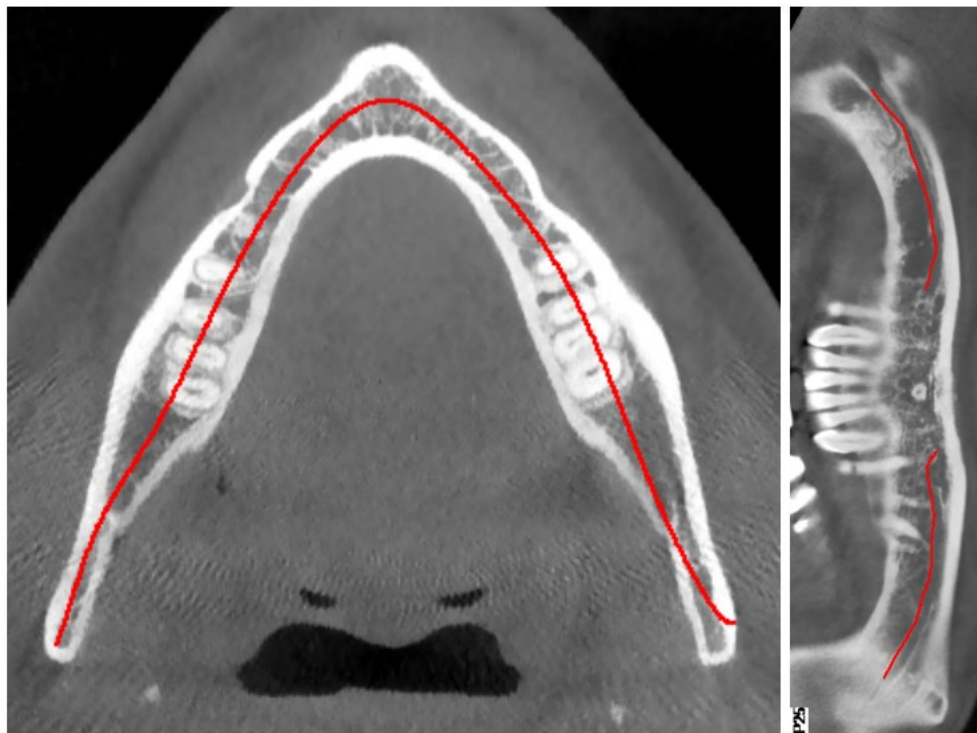
Actionable
Output

Reliable
Supervision

Dense 3D masks are the exception. Sparse labels are what clinics already produce.

01

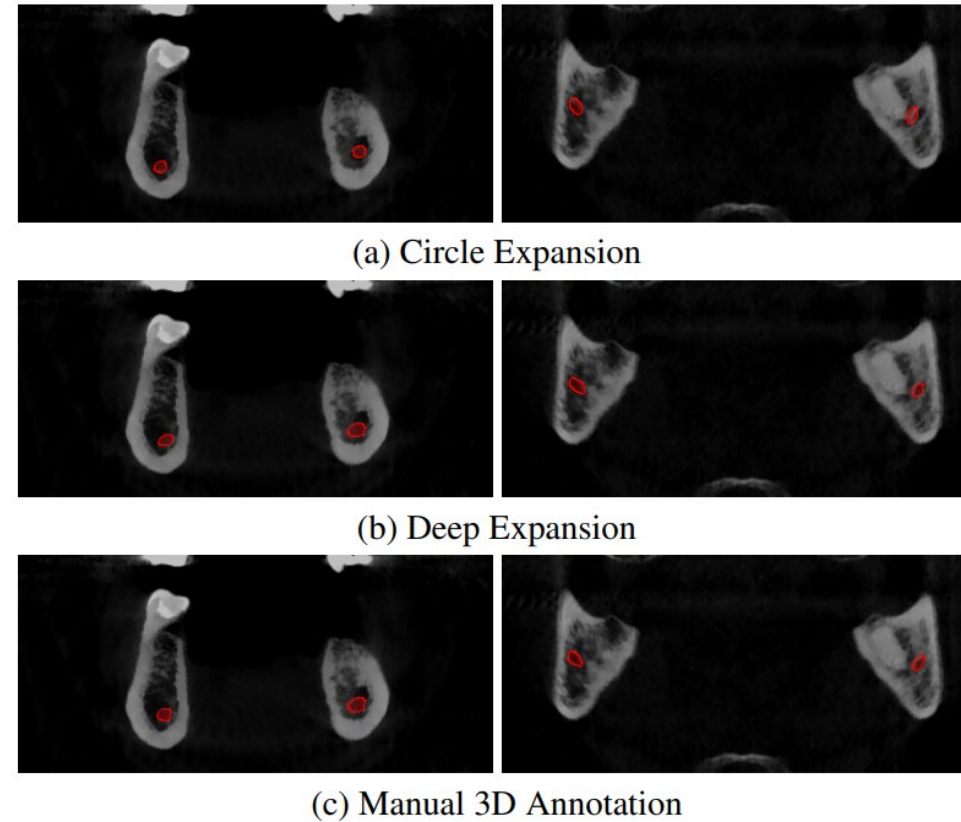
What the clinic
gives you



SPARSE BY DEFAULT

02

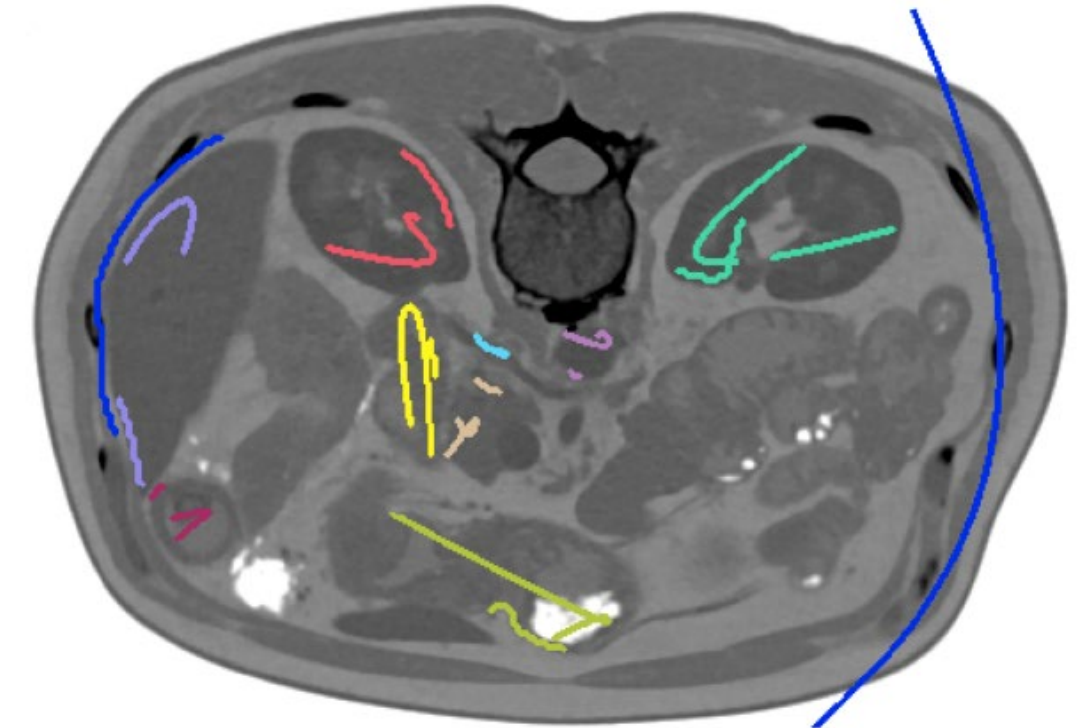
Propagate it



OUR ROUTE — CVPR 2022

03

Or ask for less



THE INTERACTIVE/SCRIBBLE ROUTE

[1] Cipriano M., et al. "Improving Segmentation of the Inferior Alveolar Nerve through Deep Label Propagation", CVPR 2022.

[2] Luo X., et al. "Scribble-Supervised Medical Image Segmentation via Dual-Branch Network and Dynamically Mixed Pseudo Labels Supervision", MICCAI 2022.

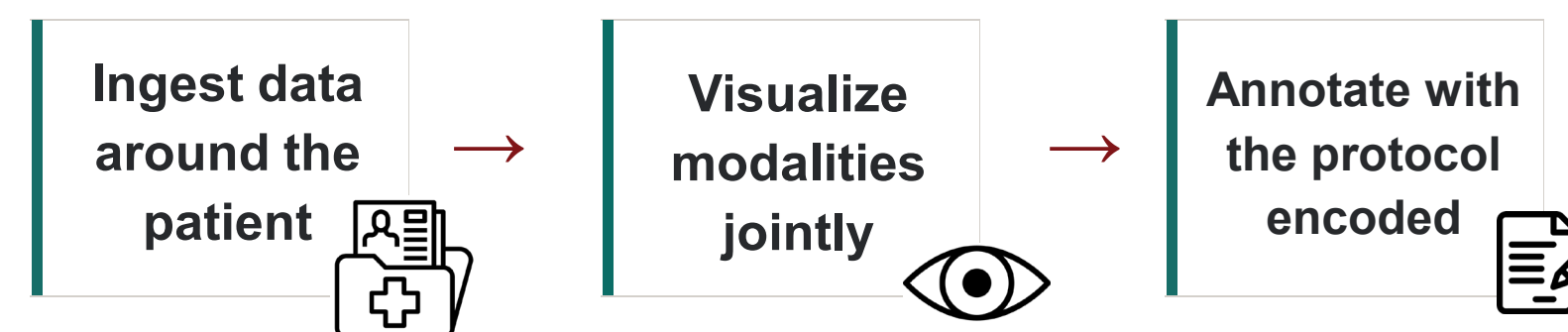
[3] Gotkowski, K. , et al. "Revisiting 3D Medical Scribble Supervision: Benchmarking Beyond Cardiac Segmentation", MICCAI 2025.



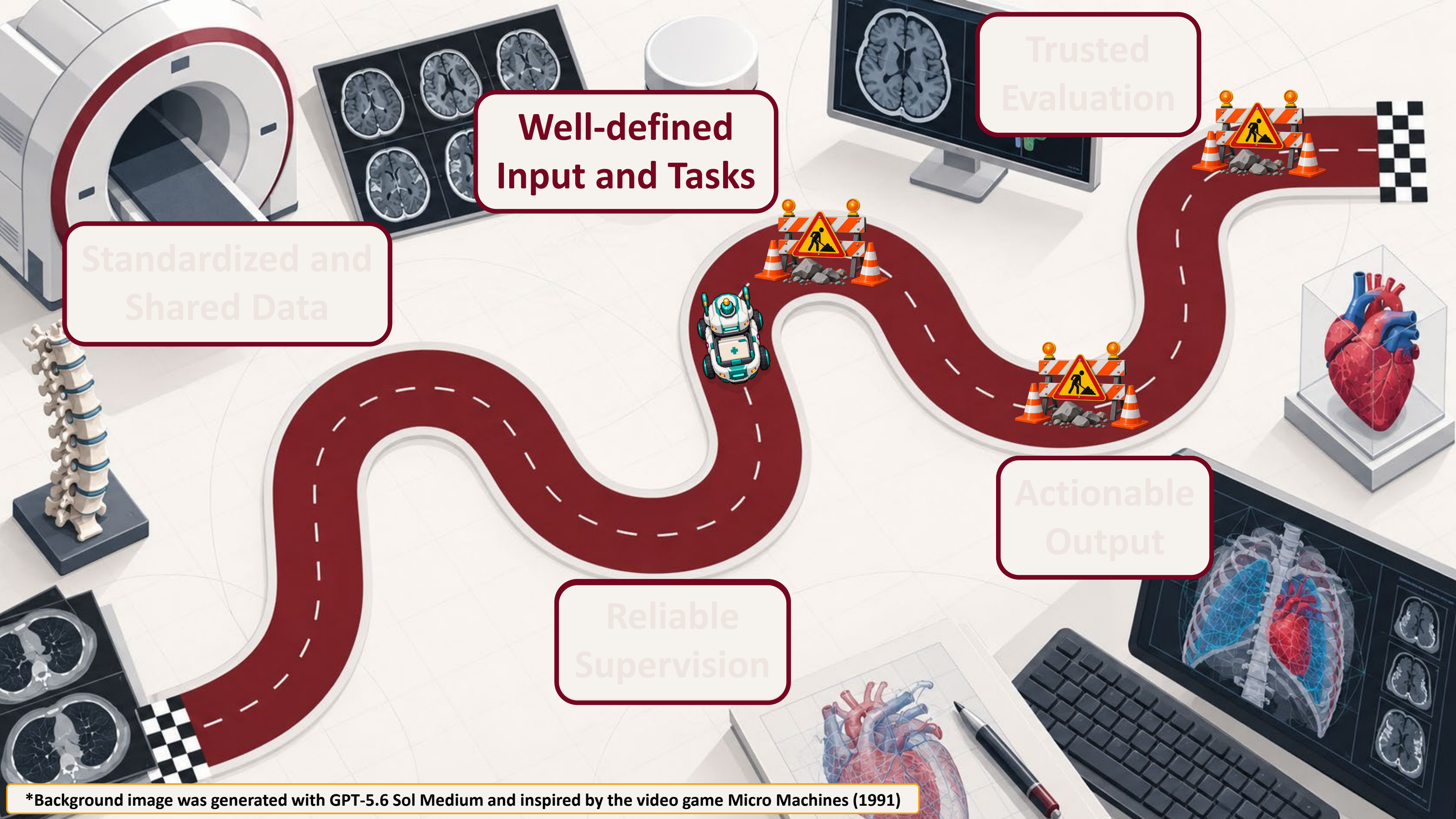
Reducing annotation time does not solve annotation consistency.

- We moved the protocol itself into the infrastructure.
- Data are organized around the patient. Data can be visualized jointly, and the annotation protocol is encoded in the platform.
- Clinicians work collaboratively, instead of producing independent datasets that we later try to reconcile.

Video not available in the PDF!



Poster Session 2 — ExHall, Poster #325, 4:30 PM CEST



**Well-defined
Input and Tasks**

**Standardized and
Shared Data**

**Trusted
Evaluation**

**Actionable
Output**

**Reliable
Supervision**

A model that requires every input to be present is a model that cannot be deployed.

WHY AN INPUT GOES MISSING

-  Time in the acquisition slot
-  Acquisition costs
-  Patient condition
-  Equipment available at that center

Accuracy should fall in proportion to the information lost — and the model should stay usable at every subset of inputs.

THE MRI EXAMPLE — BRATS

Brain tumor segmentation assumes four sequences — T1, T1c, T2, FLAIR. nnU-Net leads on the complete set, but the complete set is rarely what arrives. IM-Fuse degrades gracefully instead.

**Mamba-based
Fusion Model**
IM-Fuse

A deployed model is not a finished artifact.

WHAT CHANGES AFTER DEPLOYMENT

New anatomical structures to cover

New scanners or centres appear

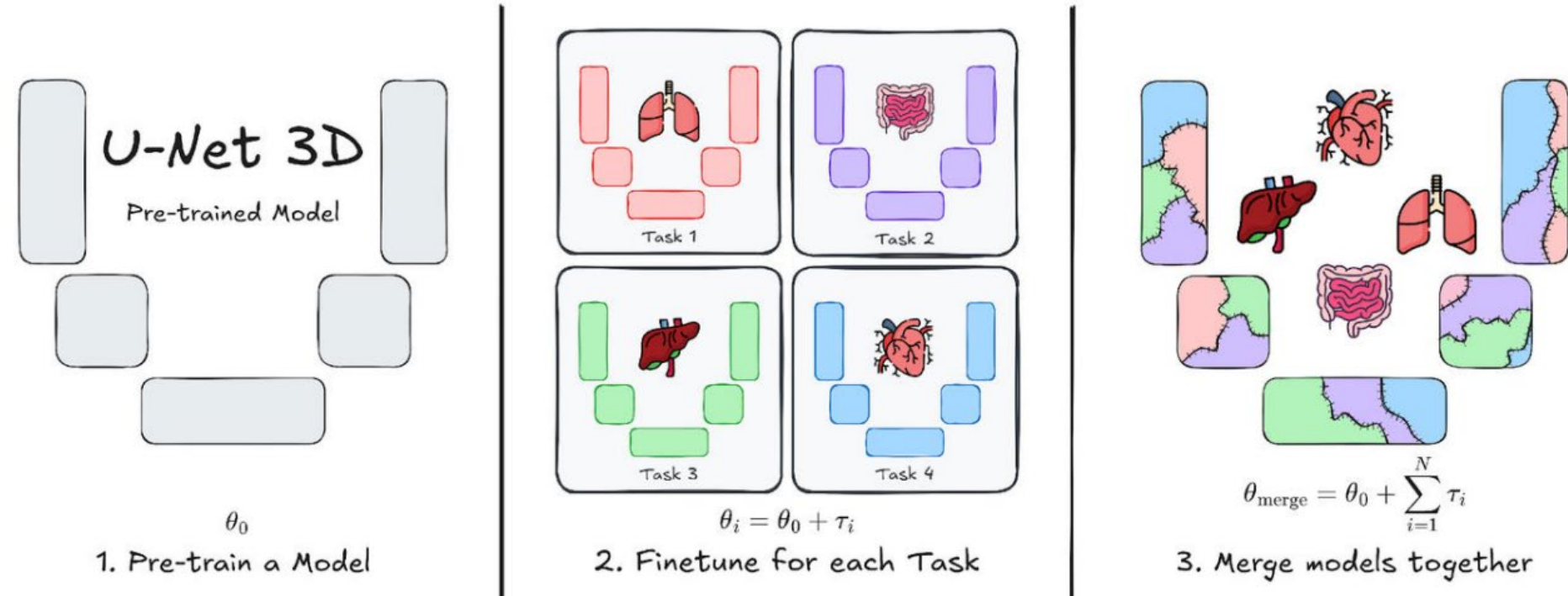
New clinical requirements emerge

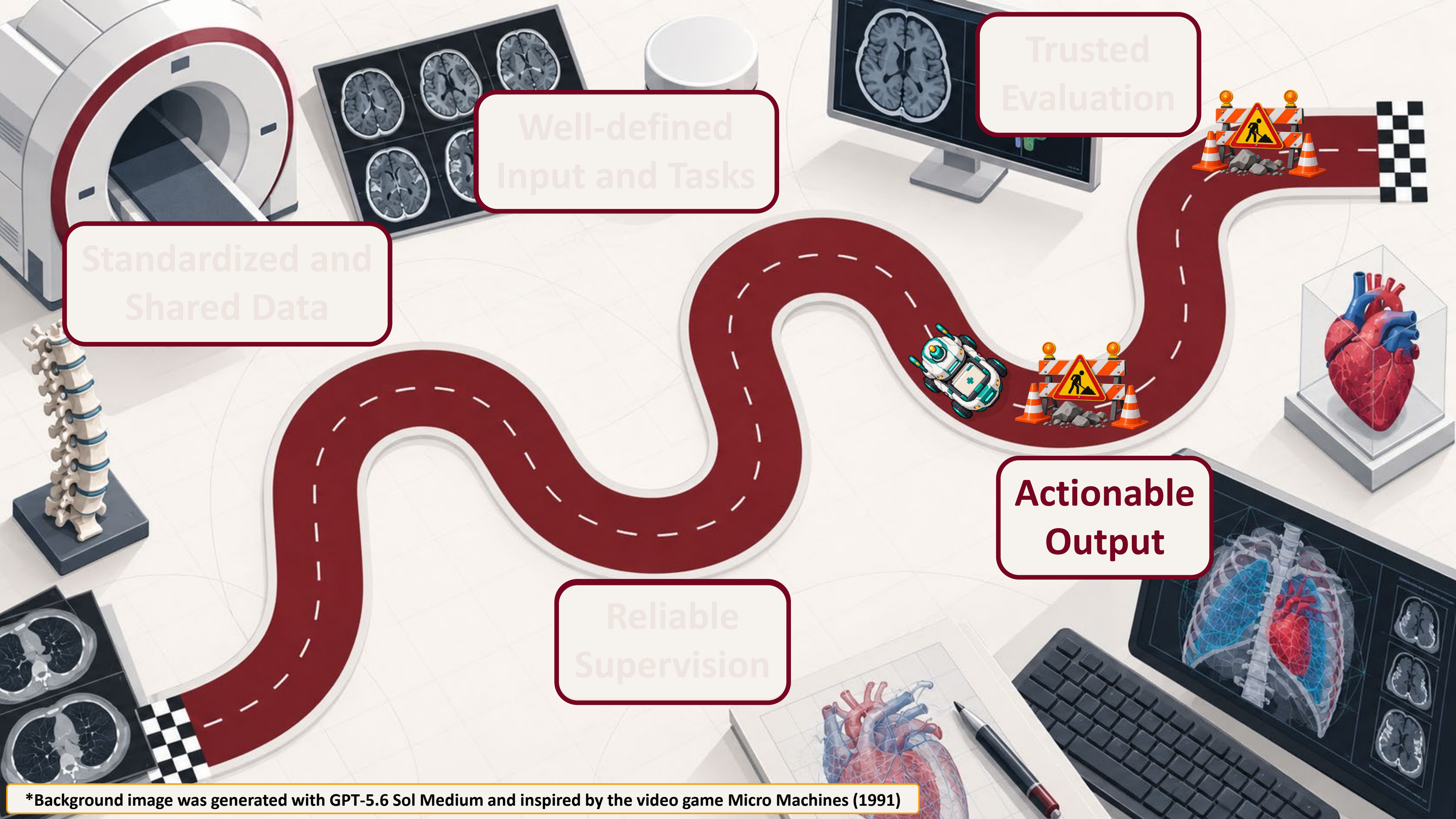
The original training data are gone

Adding a structure should not mean retraining everything from scratch.

THE CT/CBCT EXAMPLE — U-NET TRANSPLANT

Instead of one monolithic network, train specialists separately and merge them. Pre-training that encourages wide minima makes those task vectors combine cleanly — evaluated on ToothFairy2 and BTCV Abdomen.





Standardized and
Shared Data

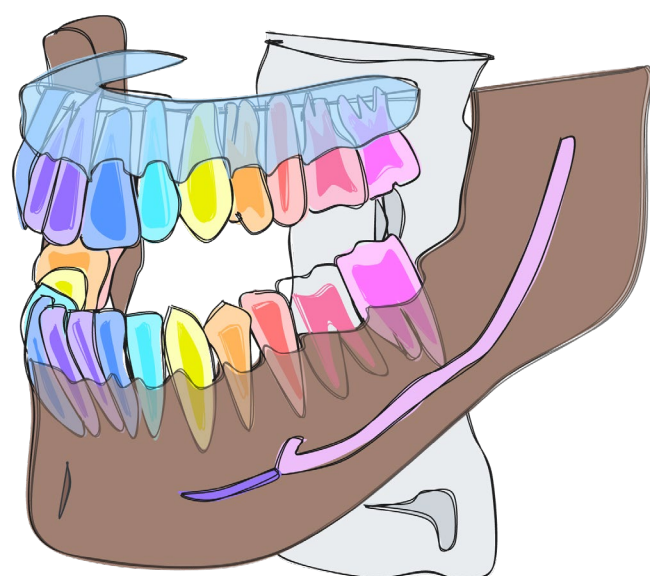
Well-defined
Input and Tasks

Trusted
Evaluation

Actionable
Output

Reliable
Supervision

A patient is not a volume



SEGMENTED VOLUME

PATH A — RAW GEOMETRY

tooth 38

1.8 cm³

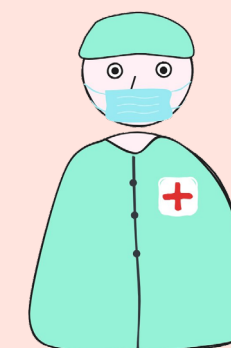
Accurate, but not yet a clinical statement



PATH B — REPORT GENERATION

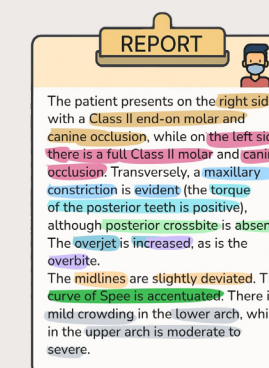
Findings stated in clinical language

Relations, severity, laterality; what the volume implies



Clinician's intervention

A human has to translate the numbers into findings



Directly actionable

The output now carries the clinical meaning

Segmentation describes anatomy. Reports communicate clinical meaning.

ToothFairy4 moves the benchmark from masks to structured findings.

WHAT WE RELEASED

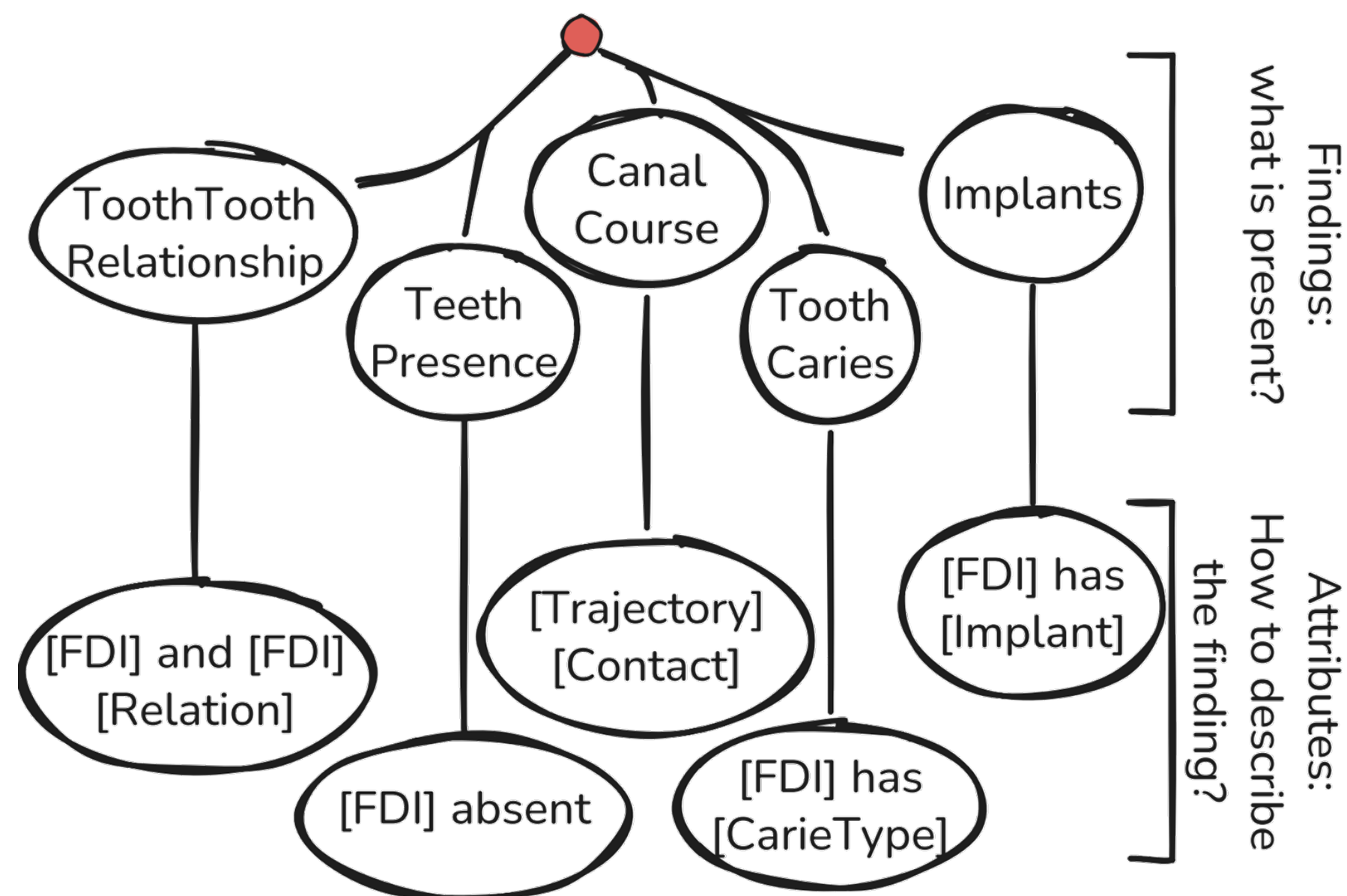
893 clinical reports for 529 public CBCT volumes

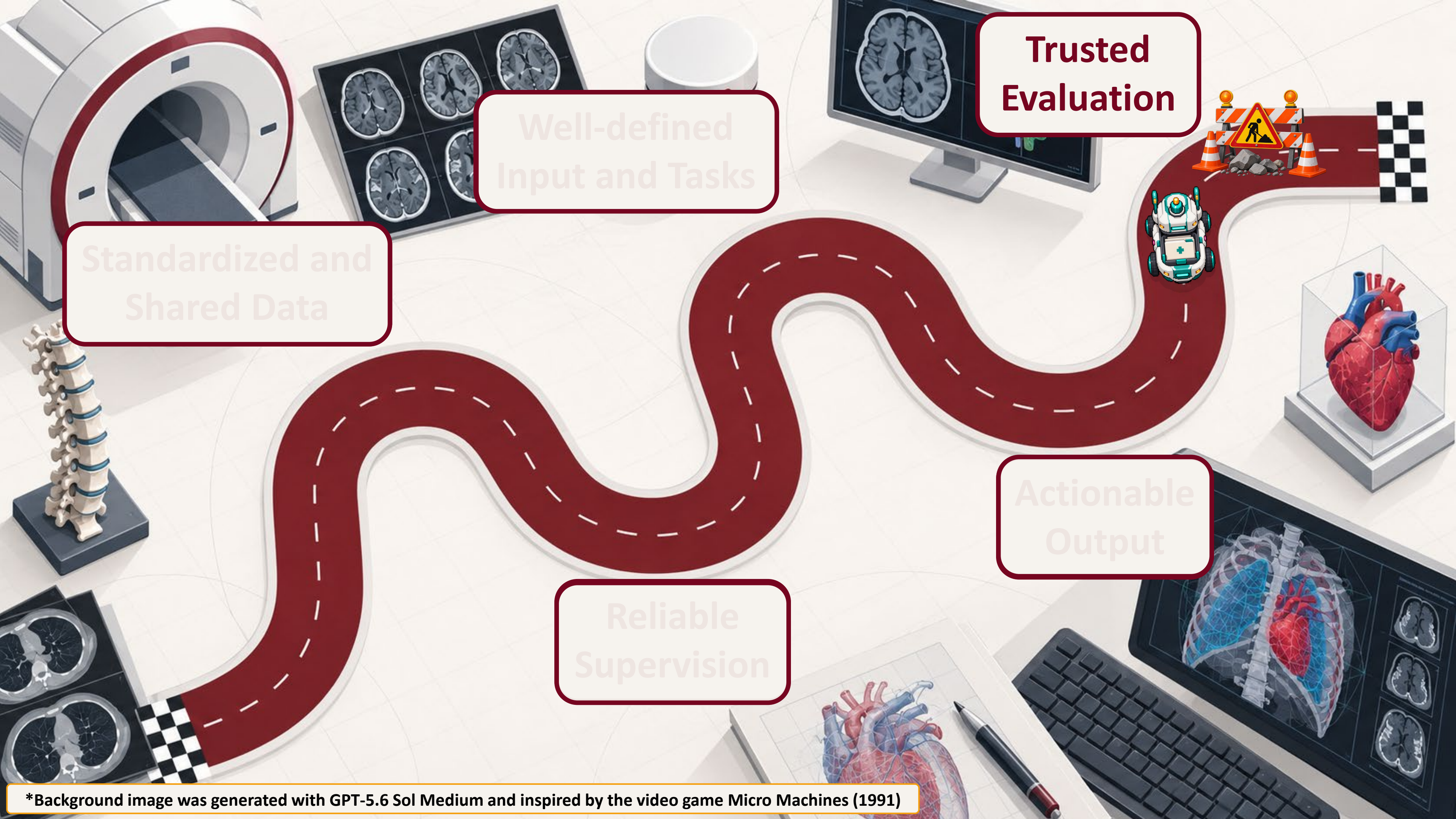
A clinician-designed OWL ontology, 13 finding types

Every report validated as an RDF/Turtle instance

Six metrics that decouple the three stages

Reporting posed as structured prediction — not free text.





Standardized and
Shared Data

Well-defined
Input and Tasks

Trusted
Evaluation

Actionable
Output

Reliable
Supervision

COMMON LEXICAL METRICS

BLEU n-gram precision

METEOR unigram alignment

ROUGE n-gram recall

All three score how a report is *phrased*.
None of them check whether the finding is semantically **true**.

REF

Teeth

18,

28,

38,

48

are

absent

GENERATED

Teeth

18,

28,

38,

48

are

endodontically treated

6 of 7 tokens identical → high n-gram overlap

the one token that carries the diagnosis

0.68

BLEU-4

reads as a strong match

0.00

FINDING AGREEMENT

wrong diagnosis!

So maybe a more structured metric that checks actual findings is needed....

RADFACT, SENTENCE LEVEL

Correctness

is each generated statement supported by the reference?

Completeness

does the report cover the reference findings?

LIMITATION

Both judgements are produced by an LLM — and inherit its failure modes.

LLM JUDGE · SENTENCE-LEVEL ENTAILMENT



PROMPT

hypothesis "Absence of tooth 36"

reference	edentulous 45	edentulous 48	edentulous 35	edentulous 16	edentulous 18
	restor. 46	restor. 47	endodontic 26	post+core 17	no entry: 36

JUDGE RESPONSE

label **entailment** → scored as correct

evidence "Edentulism of tooth 36"



never appears in the reference — invented

RADFACT STATUS

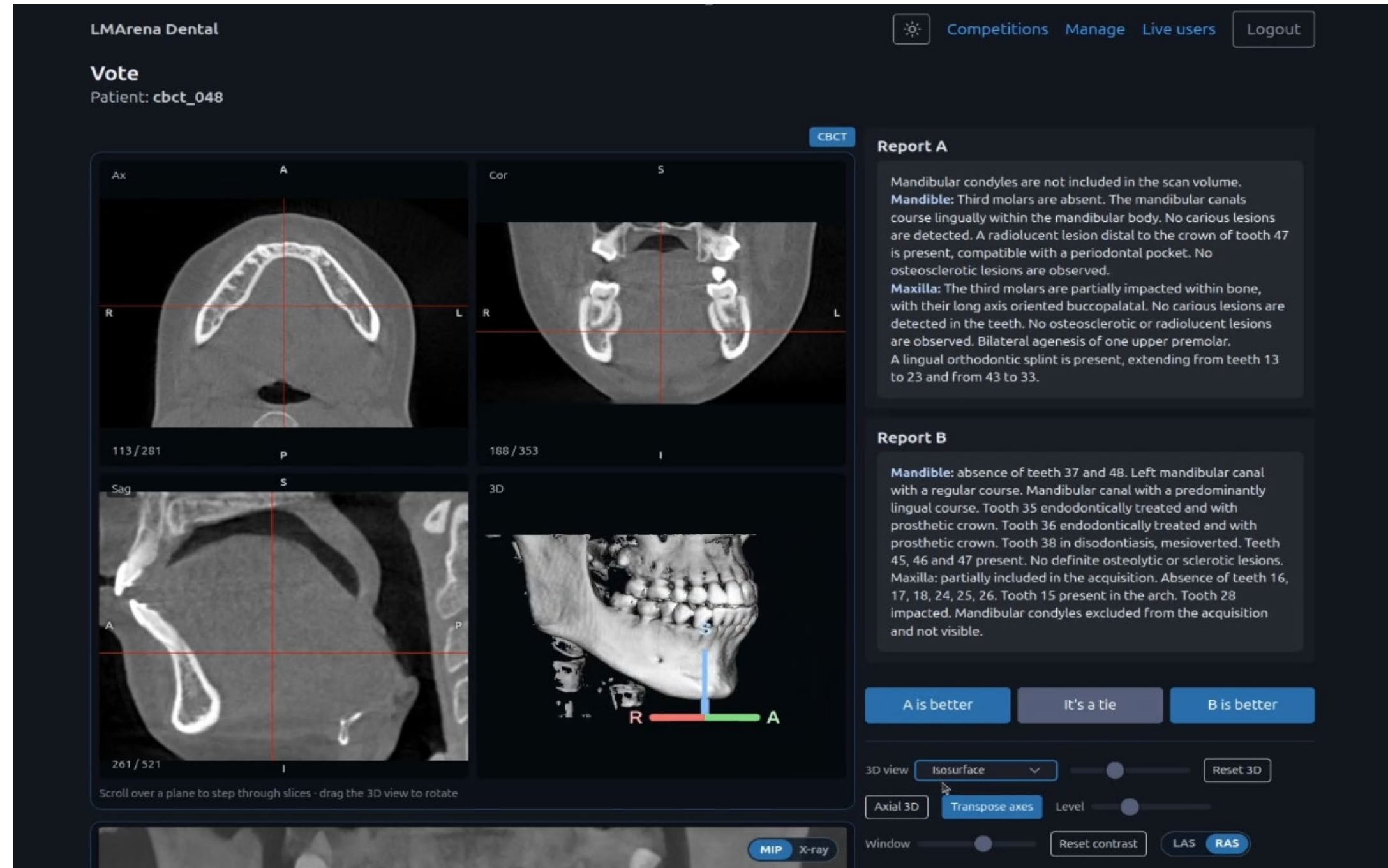
entailment

scored as correct

EVIDENCE IT CITES

Edentulism of tooth 36

never appears in the reference — invented



Video not available in the PDF!

A hospital has ten candidate models, a new scanner and no annotations.

WHAT YOU DO NOT HAVE

Labels for the target data

The source training data

Access to weights or features

A published number you can trust

Published Dice on someone else's test set does not tell you which model to deploy.

RANK BY CONSISTENCY, NOT BY SCORE

Perturb the input and see whether the prediction survives. A model whose output stays invariant is the one that actually transferred Effective Invariance, extended from classification to semantic and instance segmentation.

UNSUPERVISED

SOURCE-FREE

BLACK-BOX

No labels, no source data, no weights — only the predictions.

Estimated rankings track true target-domain performance, in 2D and 3D.

Poster Session 4 — ExHall, Poster #391, 4:00 PM CEST

PROGRESS, BUT NO PANACEA

Every direction helps — **none** of them is a panacea

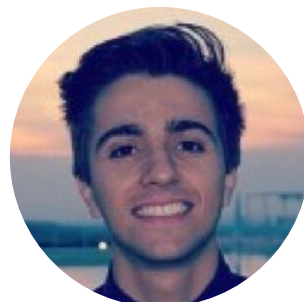
-
- 01** Shared benchmarks solved comparison, not clinical evaluation
 - 02** Human-in-the-loop can anchor clinicians to the model's bias
 - 03** Robustness is not clinic-ready, and multimodality demands more
 - 04** Clinic-ready reports need blinded reads and better metrics
-

AIMAGELAB — UNIVERSITY OF MODENA AND REGGIO EMILIA

GROUP



F. Bolelli



K. Marchesini



E. Candeloro



O. Carpentiero



N. Morelli



E. Vezzali



L. Borghi



G. Casari



G. Rosati



Y. Assefa



M. Lugli



L. Lumetti



C. Grana

CLINICAL PARTNERS

A. Anesi · M. Di Bartolomeo · A. Pellacani · P. Minafra · G. Pellacani · L. Bertoni · G. Bianchi · S. Puliatti · G. Besutti · R. Magistroni · F. Benuzzi · F. Lui · L. Lombardo · F. Cremonini

COLLABORATIONS

From Volumes to Clinical Decision

Five conditions to fill the gap between medical 3D vision and the clinic

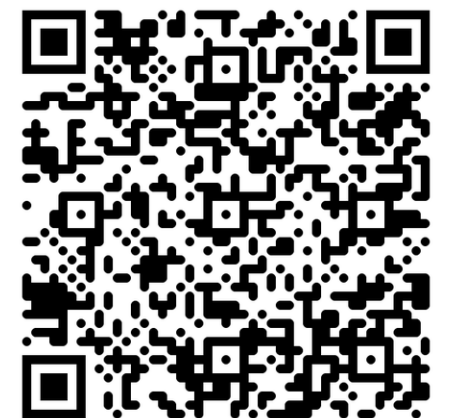
Thank You!

Federico Bolelli

AlmageLab — University of Modena and Reggio Emilia · federico.bolelli@unimore.it



UNIMORE
UNIVERSITÀ DEGLI STUDI DI
MODENA E REGGIO EMILIA



Papers, datasets, code