



Beyond Deployment: An In-House System for Continuous Quality Assurance of Deployed Auto-Segmentation Models

Jonathan Cheng, University of Pennsylvania; Nicholas Molina; Kelechi Obike

Problem & Motivation: What clinical or operational problem does your MVP address?

Deploying an AI model is not the last step. In radiation oncology, AI now drafts the organ and tumor contours that determine where radiation is delivered, but once a model is live, the harder question begins: is it still doing what the clinic needs? Deployed models require ongoing observability, and AI introduces new failure modes that traditional quality assurance was never designed to catch. We already QA our linear accelerators on a regular, rigorous schedule. We should be doing the same for our AI models, yet the standards for that don't exist yet. Before we can build them, we need to collect the data and measure how these models are performing over time. That is the real problem. We can't purchase AI tools from a vendor, plug them in, and simply trust that they keep working. Someone has to monitor them, catch the deviations, and make that performance visible to the people responsible for patient safety. Today that monitoring is largely ad-hoc, with no consistent way to measure how closely AI contours agree with the approved clinical contours, where they disagree, or how performance changes across model versions and over time. We built this to bring deployed AI models under the same standard of continuous, measurable quality assurance we already apply everywhere else in the clinic.

What You Built: Describe your MVP, tool, or workflow in plain language

We built two things in-house that work together. The first is dcm-utils, a Python package that reads radiation-oncology medical images and does the underlying geometry and metric calculations. The second is a Dashboard that turns those calculations into a picture of how our AI models are performing. The input is the DICOM data a treatment plan produces: the CT scan, the contours drawn on it, and the delivered dose. A pipeline pulls studies from our Orthanc image archive, separates clinician-approved contours from AI-generated ones using the DICOM approval status, and automatically matches structures across the different naming conventions clinicians and AI models use (for example, a clinician's "L_Lung" against a model's "Lung_L"). For each matched pair, dcm-utils computes a range of comparison metrics, including the Dice coefficient, Hausdorff distance, and surface Dice, along with per-structure shape statistics such as volume. The output is the Dashboard which contains score distributions, per-structure breakdowns, side-by-side comparisons of different AI model versions, and trends over time. It also includes a patient search, which lets us look up a specific patient's contours directly and has been invaluable for diagnosing individual contour problems and catching naming mismatches between clinical and AI structures. Every number shown traces back to a dcm-utils computation, so the Dashboard is a live, query-able view of how our deployed models are doing.

Tech Stack & Development Approach: What tools, frameworks, or methods did you use to build it?

The system has three parts plus the dcm-utils package underneath. The frontend is a React and Vite single-

page application with multiple views — clinical, model-comparison, structure-metrics, and patient-search pages — and charts built in Plotly and Recharts. The backend is a FastAPI service backed by PostgreSQL, which serves the aggregated metrics to the frontend. The data pipeline, written in Python, pulls studies from our Orthanc image archive (a free, open-source medical image server), computes the metrics, and stores them. Underneath all of it sits dcm-utils, our in-house Python package, which does the actual DICOM reading and geometry math: turning contours into 3D masks and computing spacing-aware, sub-voxel-accurate overlap and surface metrics. It handles the parts that are easy to get wrong, like anisotropic voxel spacing and coordinate-system conventions, so the pipeline can stay focused on matching, comparing, and storing. One genuinely hard problem was matching structures when clinicians and AI models name the same organ differently. We solved it in stages: a hand-curated alias list, best-overlap matching, exact name matching, and finally fuzzy matching using the Hungarian algorithm (via `scipy`), which finds the globally optimal one-to-one assignment between two sets of structure names above a similarity threshold. We built the system with Claude Code, which let us move quickly across the entire stack (database migrations, backend endpoints, and frontend charts alike) and iterate on the design as the clinical requirements sharpened.

Validation & Evidence: Have you done any informal testing or validation? If so, what did you find?

The strongest evidence is that the system runs on real clinical data. It ingests actual studies from our hospital Orthanc archive, which contains real CT scans, clinician-approved contours, and AI output. To date, our system has processed 8,097 studies across 7,422 patients. This is production evidence of function on the same data the clinic works with every day. The numbers it reports are trustworthy by construction. Every metric is a well-established segmentation measure (Dice, Hausdorff at the 95th percentile, surface Dice) computed on spacing-aware 3D masks that correctly handle anisotropic voxels and coordinate-system conventions, the details most likely to introduce silent error. Where a metric is undefined, such as an empty mask or a study with no target, the pipeline returns an explicit null rather than a misleading zero, so no fabricated value ever reaches a chart. Automated tests pin these behaviors across the pipeline, from metric storage to the API. Most importantly, the tool has already proven useful in practice. Patient search has let us diagnose individual contour problems and catch naming mismatches between clinical and AI structures that would otherwise go unnoticed, and version-to-version comparison lets us see how a newer model performs against an older one on the same anatomy. Furthermore, we can pinpoint the structures our AI models consistently struggle with and focus our efforts there. The evidence is not a one-time study; it is a system that surfaces real discrepancies on real cases, continuously.

Feedback You're Seeking: What specific feedback, help, or collaboration are you hoping to get from the CAIMI26 community?

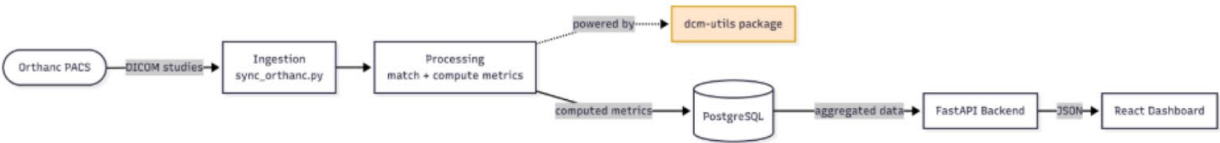
We would value your perspective on three questions. First, and most fundamental: is continuous monitoring of deployed auto-segmentation models genuinely useful, and do you see it as widely adoptable across institutions? We believe deployed AI needs ongoing quality assurance the same way a linear accelerator does, but we want to hear whether practitioners see this as a real need or a solution in search of a problem. Second, what agreement thresholds should separate an acceptable AI contour from one that genuinely needs editing, and should those thresholds differ from one structure to another? A Dice of 0.85 may be excellent for a large organ and unacceptable for a small structure near the cord, and we have no consensus on where those lines belong. Third, beyond the metrics already on our roadmap, what features or measurements would take this system to the next level? We have identified the obvious next steps, such as dose-weighted comparison, but we are most interested in the non-obvious ones: the views, signals, or capabilities that a daily user would want but that we have not thought to build.

Known Limitations & Honest Failures: What isn't working yet, or what have you already tried that failed?

Each of our main limitations comes down to the same gap: the numbers are computed correctly, but reading them still requires clinical context. The first and deepest is our definition of ground truth. We treat any contour with a DICOM approval status of APPROVED as the clinical reference, but physicians sometimes approve an AI contour with little or no editing even when it does not truly qualify as ground truth. When that happens, an AI-versus-clinical comparison is really AI-versus-approved-AI, which inflates agreement and overstates model quality. Such a contour shows near-perfect Dice and near-zero Added Path Length, so the metrics we are adding may themselves become a way to flag these cases. The second is clinical convention. When contouring breast, clinicians often crop several millimeters from the skin, a step the AI model does not know to take. Our system flags this as an AI-versus-clinical disagreement even though it is not a model error. This is the kind of disagreement that convention-aware or distance-weighted comparison is meant to handle, and addressing it is on our roadmap. The third is structure naming. Our multi-stage matcher has made false matches rare, but no matcher can fully compensate for inconsistent naming at the source. A physician focused on one breast may label it simply Breast rather than Breast_L or Breast_R, leaving no way to tell which side it is and risking a match against the wrong side. This is why reliable naming, and the standards behind it, remains an open problem.

Potential for Impact: If this MVP were fully developed and deployed, who would benefit and how?

The immediate beneficiary is the whole department, not one role. Our system gives physicists and dosimetrists a shared, current view of auto-segmentation quality by structure and model version, replacing ad-hoc spot checks with a consistent record. It shows physicians where a model is systematically weak, so review attention goes where it is most needed, which is a patient-safety benefit as much as an efficiency one. It gives quality and safety an auditable, continuous account of how a deployed model is performing over time. And it gives administrative leadership objective data for decisions that are otherwise made on vendor claims alone: whether to adopt a model, keep it, or replace it. That last point is not hypothetical. When a newer model version was introduced, our system showed that agreement with approved clinical contours did not improve, and on certain structures it declined. Without continuous monitoring, that regression would likely have gone unnoticed, and the newer version would have been adopted on the assumption that newer means better. This is exactly the failure mode that motivated the project: a deployed model can quietly get worse, and only measurement will tell you. The broadest impact reaches past our own department. Standards for monitoring deployed AI models do not largely exist yet and building them will require real-world performance data collected consistently over time. That is precisely what our system produces, at the scale of thousands of patients. We see it as an early contribution toward the evidence base such standards will need.



RAOINC.AI AUTOSEG

Clinical Segmentation Metrics

Clinical AI Dashboard

Overview of all clinical segmentation metrics
Charts include all AI model versions (2.0 and 5.0) compared against clinical data. For model-specific breakdowns, see the Model Comparison page.

STUDY DATE

From To

OVERALL AVERAGE DICE SCORE

0.92

AVG SURFACE DISTANCE (MM)

1.69

ROBUST HUISBOFF (YX)

9.47

SURFACE SIZE (MM)

0.88

PATIENTS PROCESSED

7161

ROIS ANALYZED

790

TOTAL PREDICTIONS ANALYZED

107203

DATE RANGE

2021-04-01 - 2026-06-22

< >

Cardiovascular Structures

5 / 13

Cardiovascular Structures

DICE Score

DICE Score Distribution

A.CircumflexArtery (n=104)

A.CircumflexVein (n=105)

A.Coronary_L (n=107)

A.LAB (n=284)

A.Pulmonary (n=35)

Aorta (n=107)

Aorta_The_Asc (n=78)

Aorta_The_Des (n=322)

Atrial_L (n=11)

Atrial_R (n=11)

Coronary_Sinus (n=23)

Bicuspid (n=3)

Heart (n=4089)

Ventric_L (n=103)

Ventric_L_Ant (n=17)

Ventric_L_Pst (n=8)

Ventric_L_Ant (n=17)

Ventric_L_Pst (n=8)

Ventric_L_Ant (n=17)

Ventric_L_Pst (n=8)

Ventric_L_Ant (n=17)

Ventric_L_Pst (n=8)

The screenshot displays the MONAI AI Lab interface, specifically the 'Model Performance Metrics' section for the 'Abdominal (Non-Pelvic) Structures' task. The interface is divided into several panels:

- Top Panel:** Shows the model name 'MONAI 3.0.0-19.3.0' and the task 'Abdominal (Non-Pelvic) Structures'. It includes a 'Model Performance Metrics' button.
- Table:** A table comparing the performance of the 'MONAI 3.0.0-19.3.0' model against the 'MONAI 3.0.0-19.3.0' baseline. The table has columns for 'Model', 'Dice Score', 'Dice Score (std)', 'Dice Score (min)', 'Dice Score (max)', and 'Dice Score (mean)'. The 'MONAI 3.0.0-19.3.0' model shows a Dice Score of 0.901, which is higher than the baseline of 0.899.
- Bottom Panel:** A 'DICE Score Distribution' plot showing the distribution of Dice scores for various structures. The x-axis lists structures: Pan, Spleen, Liver, Stomach, Gallbladder, Kidney, Adrenal, Spleen, Liver, Stomach, Gallbladder, Kidney, Adrenal, Spleen, Liver, Stomach, Gallbladder, Kidney, Adrenal. The y-axis represents the 'Dice Score' from 0.0 to 1.0. The plot shows that the 'MONAI 3.0.0-19.3.0' model (red) has a higher Dice score distribution compared to the baseline (blue).

Patient Search

Look up a single patient by ID to view their stock history and compare their result to each of our validated models.
For more information on the data model and use of the results, see our [FAQ](#) or [Contact Us](#) page.

SEARCH

SEARCH

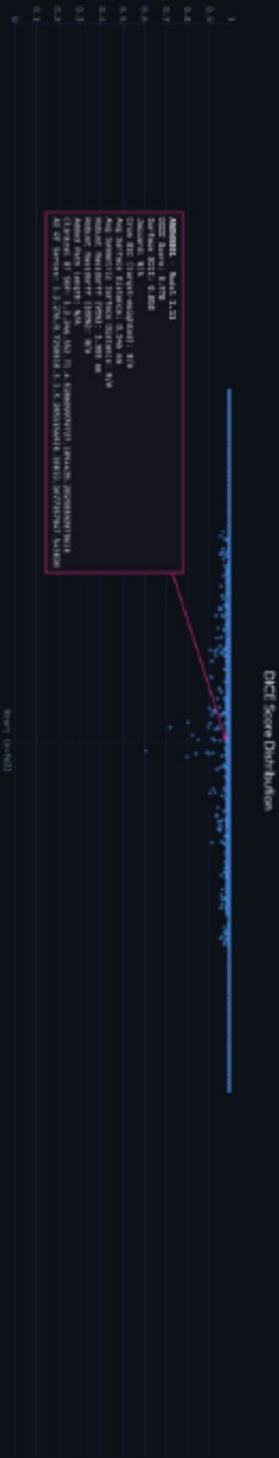
RESULTS FOR PATIENT: ASB00001

Model 1, 12

Model 1 R ² : 0.73 AUC: 0.86 AUPRC: 0.87	Model 2 R ² : 0.73 AUC: 0.86 AUPRC: 0.87	Model 3, 24 R ² : 0.84 AUC: 0.91 AUPRC: 0.92	Model 4 R ² : 0.89 AUC: 0.92 AUPRC: 0.93	Model 5 R ² : 0.90 AUC: 0.93 AUPRC: 0.94	Model 6 R ² : 0.90 AUC: 0.93 AUPRC: 0.94	Model 7 R ² : 0.91 AUC: 0.94 AUPRC: 0.95
---	---	---	---	---	---	---

Model 8, 24
R²: 0.91
AUC: 0.94
AUPRC: 0.95

Model 9



DATA STRUCTURE: DCE, COMPASSION

For more information on the data model and use of the results, see our [FAQ](#) or [Contact Us](#) page.

Model 1: 0.73, Model 2: 0.73, Model 3: 0.84, Model 4: 0.89, Model 5: 0.90, Model 6: 0.90, Model 7: 0.91, Model 8: 0.91, Model 9: 0.91, Model 10: 0.91, Model 11: 0.91, Model 12: 0.91, Model 13: 0.91, Model 14: 0.91, Model 15: 0.91, Model 16: 0.91, Model 17: 0.91, Model 18: 0.91, Model 19: 0.91, Model 20: 0.91, Model 21: 0.91, Model 22: 0.91, Model 23: 0.91, Model 24: 0.91, Model 25: 0.91, Model 26: 0.91, Model 27: 0.91, Model 28: 0.91, Model 29: 0.91, Model 30: 0.91, Model 31: 0.91, Model 32: 0.91, Model 33: 0.91, Model 34: 0.91, Model 35: 0.91, Model 36: 0.91, Model 37: 0.91, Model 38: 0.91, Model 39: 0.91, Model 40: 0.91, Model 41: 0.91, Model 42: 0.91, Model 43: 0.91, Model 44: 0.91, Model 45: 0.91, Model 46: 0.91, Model 47: 0.91, Model 48: 0.91, Model 49: 0.91, Model 50: 0.91, Model 51: 0.91, Model 52: 0.91, Model 53: 0.91, Model 54: 0.91, Model 55: 0.91, Model 56: 0.91, Model 57: 0.91, Model 58: 0.91, Model 59: 0.91, Model 60: 0.91, Model 61: 0.91, Model 62: 0.91, Model 63: 0.91, Model 64: 0.91, Model 65: 0.91, Model 66: 0.91, Model 67: 0.91, Model 68: 0.91, Model 69: 0.91, Model 70: 0.91, Model 71: 0.91, Model 72: 0.91, Model 73: 0.91, Model 74: 0.91, Model 75: 0.91, Model 76: 0.91, Model 77: 0.91, Model 78: 0.91, Model 79: 0.91, Model 80: 0.91, Model 81: 0.91, Model 82: 0.91, Model 83: 0.91, Model 84: 0.91, Model 85: 0.91, Model 86: 0.91, Model 87: 0.91, Model 88: 0.91, Model 89: 0.91, Model 90: 0.91, Model 91: 0.91, Model 92: 0.91, Model 93: 0.91, Model 94: 0.91, Model 95: 0.91, Model 96: 0.91, Model 97: 0.91, Model 98: 0.91, Model 99: 0.91, Model 100: 0.91

Model 1: 0.73

Model 2: 0.73

Model 3: 0.84

Model 4: 0.89

