

AI Driven MedFusion -X: An Extended Multimodal CNN–Transformer Framework for Biomedical Image and Clinical Data Fusion

Dr. Kommu Naveen^{1*}, Dr. Sunil Vijaya Kumar Gaddam², Mediga Naga Seshudu³, Dr. M P Mallesh⁴, T. Laxmi Prasanna⁵, Dr. Madhukar. G⁶, Dr. Bhaludra R Nadh Singh⁷

¹Professor, Department of Computer Science Engineering- AIML, Kasireddy Narayanareddy College of Engineering and Research, Abdullapurmet, Near Ramoji Film City, Ranga Reddy District, Hyderabad- 501513, Telangana State, India.

naveenkarunya@gmail.com

²Professor & Dean, Department of Computer Science Engineering, Rajeev Gandhi Memorial College of Engineering & Technology (Autonomous), Nandyal - 518 501, Andhra Pradesh. India. sunilvk@gmail.com

³Ph.D Research Scholar, Department of Computer Science Engineering, JNT University, Ananthapur - 515002, Andhra Pradesh, India. nagaseshudu1000@gmail.com

⁴Associate Professor, Department of Mathematics, Koneru Lakshmaiah Education Foundation, Aziznagar, Hyderabad, Telangana, India - 500075. malleshmardanpally@gmail.com

⁵Assistant Professor, Department of Computer Science Engineering – CSE (DS) , TKR College of Engineering and Technology, Affiliated to JNTU Hyderabad, TKRCET, Medbowli, Meerpet, Saroornagar, Hyderabad -500097, prasannabejugam92@gmail.com

⁶Associate Professor, Computer Science and Engineering, School of Engineering and Technology, Kaveri University, Gowraram Village, Wargal Mandal, Siddipet District, Telangana, 502279, India. madhu.mani536@gmail.com

⁷Professor, Computer Science Engineering, Neil Gogte Institute of Technology, Hyderabad, Telangana, India-500 088. brn.singh@gmail.com

***Corresponding Author**

Email: naveenkarunya@gmail.com

Received: 25th May, 2026; **Revised:** 6th June, 2026; **Accepted:** 8th June, 2026; **Available Online:** 09th June, 2026

ABSTRACT

The rapid growth of artificial intelligence (AI), deep learning (DL), and multimodal representation learning has created new opportunities for developing intelligent biomedical image-analysis systems capable of integrating heterogeneous clinical evidence. Conventional convolutional neural networks (CNNs) have demonstrated strong capabilities in extracting local spatial and hierarchical visual features from biomedical images, whereas Transformer architectures provide an effective mechanism for modeling long-range contextual dependencies. However, many existing biomedical AI systems remain predominantly image-centric, task-specific, and dependent on large quantities of expert-annotated data. Such limitations can adversely affect their generalizability across medical institutions, imaging devices, acquisition protocols, patient populations, and clinical environments. Furthermore, clinical decision-making frequently requires integration of imaging findings with demographic, physiological, laboratory, electronic health record (EHR), and textual information. This paper proposes **MedFusion-X**, a multimodal CNN–Transformer framework designed for biomedical image and clinical-data fusion. The proposed architecture integrates CNN-based local feature extraction, Transformer-based global contextual representation, clinical-data encoding, cross-modal attention, adaptive feature fusion, self-supervised representation learning, uncertainty estimation, and explainable artificial intelligence (XAI). The framework is designed to learn complementary relationships between biomedical images and structured clinical information rather than treating these modalities as independent sources. The proposed workflow consists of patient-level data curation, modality harmonization, quality control, preprocessing, self-supervised pretraining, multimodal representation learning, cross-modal attention, adaptive fusion, task-specific prediction, calibration, uncertainty estimation, and explainability analysis. MedFusion-X is intended to support downstream biomedical applications including disease classification, lesion detection, image segmentation, disease grading, prognosis, and clinical risk prediction. The evaluation methodology incorporates accuracy, precision, sensitivity/recall, specificity, F1-score, ROC-AUC, PR-AUC, calibration error, computational efficiency, robustness under distribution shift, and external validation. For segmentation applications, Dice similarity coefficient, Jaccard index, Hausdorff distance, and sensitivity are additionally recommended. The proposed experimental design includes ablation studies to independently assess the contribution of CNN representations, Transformer representations, self-supervised learning, clinical fusion, cross-modal attention, and adaptive fusion. Explainability is evaluated using image-attribution and clinical-feature-importance methods, while uncertainty analysis is used to quantify prediction reliability.

The proposed MedFusion-X framework provides a unified research direction for developing transferable and interpretable multimodal biomedical AI systems. Importantly, numerical performance values and clinical superiority claims must be established through reproducible experiments using independent datasets and should not be inferred solely from the proposed architecture. The framework therefore provides a methodological foundation for investigating whether multimodal CNN–Transformer representation learning can improve predictive performance, generalization, interpretability, calibration, and clinical usefulness compared with conventional unimodal biomedical image-analysis models. This extended study presents

MedFusion-X, a multimodal CNN–Transformer framework for biomedical image and clinical data fusion. The framework addresses limitations of conventional image-centric medical AI, including dependence on expert annotations, weak integration of structured clinical information, limited long-range modeling, domain shift, incomplete modalities, calibration deficiencies, and limited interpretability. MedFusion-X combines CNN local feature extraction, Transformer global contextual modeling, clinical-data encoding, cross-modal attention, adaptive modality weighting, self-supervised representation learning, uncertainty estimation, and explainable artificial intelligence (XAI). The extended methodology introduces patient-level data partitioning, external validation, missing-modality analysis, robustness testing, subgroup evaluation, calibration assessment, statistical confidence intervals, ablation experiments, computational profiling, and reproducibility controls. The framework supports classification, lesion detection, segmentation, grading, prognosis, and risk-oriented downstream tasks. Numerical performance values and statistical claims must be generated from actual experiments and are not fabricated in this draft. The study design is intended to provide a rigorous basis for evaluating whether multimodal CNN–Transformer learning can improve generalization, label efficiency, interpretability, reliability, and clinical relevance over conventional biomedical image-analysis methods.

Keywords: Multimodal Medical AI; Biomedical Image Analysis; CNN; Vision Transformer; Clinical Data Fusion; Self-Supervised Learning; Cross-Modal Attention; Explainable AI; Uncertainty Estimation; Medical Imaging; Deep Learning; Clinical Decision Support.

How to cite this article: Kommu Naveen, Sunil Vijaya Kumar Gaddam, Mediga Naga Seshudu, M P Mallesh, T. Laxmi Prasanna, Madhukar. G, Bhaludra R Nadh Singh. AI Driven MedFusion -X: An Extended Multimodal CNN–Transformer Framework for Biomedical Image and Clinical Data Fusion. *Int J Drug Deliv Technol.* 2026;16(79s):636-638. DOI: 10.25258/ijddt.16.79s.118.

Source of support: Nil.

Conflict of interest: Nil.

INTRODUCTION

Biomedical imaging provides critical information for diagnosis, screening, staging, treatment planning, and longitudinal monitoring. MRI, CT, PET, X-ray, ultrasound, digital pathology, and related modalities generate high-dimensional observations containing anatomical, physiological, and pathological information. Deep learning has enabled automated extraction of discriminative representations, but many systems remain optimized for a single task and a narrowly defined dataset. CNNs provide strong local feature extraction and hierarchical spatial representation. However, abnormalities can involve relationships between spatially separated structures, multiple sequences, or global anatomical context. Transformer architectures address part of this limitation through self-attention, allowing image tokens to interact over long spatial ranges. A hybrid CNN–Transformer design can therefore combine local inductive bias with global contextual modeling. Clinical decisions rarely depend on imaging alone. Age, sex, symptoms, laboratory measurements, history, pathology, medication, radiology reports, and EHR information can provide complementary evidence. MedFusion-X therefore treats image and clinical information as interacting modalities rather than independent inputs. Biomedical imaging has become an essential component of modern healthcare because medical images provide non-invasive information about anatomical structures, tissue characteristics, pathological changes, and disease progression. Modalities such as magnetic resonance imaging (MRI), computed tomography (CT), positron emission tomography (PET), digital pathology, ultrasound, mammography, and X-ray imaging are routinely used for diagnosis, screening, staging, treatment planning, and monitoring. However, biomedical images are inherently high-dimensional and frequently contain complex spatial patterns that may be difficult to identify consistently through manual examination alone.

The development of AI and DL has substantially changed biomedical image analysis. CNNs have demonstrated remarkable capability in learning hierarchical visual representations directly from raw or minimally processed medical images. Through successive convolutional layers, CNNs can identify edges, textures, shapes, anatomical boundaries, lesion characteristics, and higher-order pathological patterns. CNN-based architectures have consequently been adopted for classification, segmentation, detection, registration, reconstruction, and image synthesis. Nevertheless, CNNs possess an inherent locality bias. Although deep convolutional layers can obtain increasingly large receptive fields, modeling long-range relationships explicitly can remain challenging. In medical imaging, this is important because the interpretation of a pathological region may depend on spatial relationships with distant anatomical structures or multiple regions within an image or volume. Transformer architectures provide an alternative representation-learning mechanism based on self-attention. Unlike conventional convolutional operations, self-attention enables the model to establish relationships between distant image tokens. Vision Transformers and hybrid CNN–Transformer architectures have therefore become increasingly important in medical image analysis.

However, the adoption of Transformers does not completely solve the challenges of biomedical AI. Medical datasets are often limited, heterogeneous, highly imbalanced, institution-specific, and expensive to annotate. A model trained using data from one hospital may encounter substantial domain shift when transferred to another hospital. Scanner manufacturer, field strength, acquisition protocol, reconstruction algorithm, image resolution, preprocessing, demographic characteristics, disease prevalence, and clinical workflow can all influence the data distribution. Another fundamental limitation is the separation between medical imaging and clinical information. In actual clinical decision-making,

physicians rarely interpret an image in complete isolation. Imaging findings may be combined with age, sex, symptoms, medical history, laboratory measurements, medications, pathology results, treatment history, radiology reports, and other EHR information. Consequently, an image-only AI system represents only one component of the available clinical evidence. Multimodal learning provides a mechanism for addressing this limitation. By jointly modeling biomedical images and clinical information, a multimodal system can potentially learn complementary relationships between visual and non-visual evidence. However, effective multimodal learning is more complex than simply concatenating feature vectors because image and clinical data possess substantially different structures, dimensionalities, noise characteristics, and missing-data patterns. The proposed **MedFusion-X** framework addresses these challenges through a multimodal CNN–Transformer architecture that combines local visual representation learning, global contextual modeling, clinical-data encoding, cross-modal attention, adaptive feature fusion, self-supervised learning, uncertainty estimation, and explainability. The central hypothesis is that jointly learned image–clinical representations can provide more transferable and clinically meaningful information than conventional image-only models, particularly when evaluated under limited-label, external-domain, and distribution-shift conditions.

2. Motivation

- Reduce dependence on expensive expert-labeled datasets through self-supervised pretraining.
- Combine local CNN representations with global Transformer representations.
- Integrate heterogeneous clinical information through a dedicated clinical encoder.
- Learn image–clinical interactions using cross-modal attention rather than simple concatenation.
- Improve robustness to institution, scanner, acquisition, and population shifts.
- Quantify predictive uncertainty and calibration rather than reporting accuracy alone.
- Provide image-region and clinical-feature explanations and evaluate their faithfulness.
- Support missing-modality and incomplete-clinical-data scenarios.
- Create a reproducible framework adaptable to multiple biomedical tasks.

3. Problem Statement

Let I_i denote the biomedical image or image volume of patient i and C_i denote associated clinical information. The objective is to learn a multimodal representation $Z_i = F_\theta(I_i, C_i)$ that preserves clinically relevant visual and non-visual information and supports downstream prediction while remaining robust to domain variation. The central research question is whether a jointly learned image–clinical representation can improve discrimination, generalization, calibration, interpretability, and reliability compared with unimodal

and simple-fusion baselines.

Let a biomedical examination be represented by an image or image volume:

$$I_i \in \mathbb{R}^{M \times H \times W \times D}$$

where (M) denotes imaging modalities or channels and (H, W, D) represent spatial dimensions.

Let associated clinical information be represented by:

$$C_i = \{c_{i1}, c_{i2}, \dots, c_{ip}\}.$$

The objective is to learn a multimodal representation:

$$Z_i = F_\theta(I_i, C_i)$$

such that (Z_i) is informative for downstream clinical tasks while remaining robust to domain variation.

The prediction function is defined as:

$$\hat{y}_i = H_{\psi}(Z_i)$$

where (H_{ψ}) represents a task-specific prediction head.

4. Research Hypotheses

- H1: Self-supervised pretraining improves downstream performance and label efficiency.
- H2: CNN–Transformer hybrid representations outperform CNN-only and Transformer-only models under matched conditions.
- H3: Multimodal image–clinical fusion improves performance over image-only models.
- H4: Cross-modal attention improves representation quality over simple feature concatenation.
- H5: Adaptive fusion improves robustness when modality informativeness varies between patients.
- H6: External-domain performance degradation is reduced by robust multimodal representation learning.
- H7: Integrated XAI produces stable and clinically meaningful explanations.
- H8: Uncertainty estimation identifies low-confidence cases more reliably than raw confidence.

5. Main Contributions

- The major contributions of MedFusion-X are in the implementation of the research:
- A multimodal CNN–Transformer architecture for jointly representing biomedical images and clinical information.
- A hybrid visual encoder combining CNN local representations with Transformer global contextual representations.
- A cross-modal attention mechanism for learning interactions between image-derived and clinical representations.
- An adaptive feature-fusion mechanism that dynamically determines the contribution of individual modalities.

- A self-supervised pretraining strategy designed to reduce dependence on large-scale expert annotation.
- An explainable AI module integrating image-region attribution and clinical-feature importance.
- An uncertainty and calibration framework for evaluating prediction reliability.
- A comprehensive evaluation strategy incorporating internal validation, external validation, distribution-shift testing, missing-modality analysis, and ablation studies.
- A reproducible framework intended for adaptation to multiple biomedical image-analysis tasks.
- A unified multimodal CNN–Transformer architecture for biomedical image and clinical-data fusion.
- A self-supervised pretraining stage for exploiting unlabeled biomedical images.
- A cross-modal attention mechanism for learning image–clinical interactions.
- An adaptive fusion module that learns patient-dependent modality contributions.
- An uncertainty-aware and calibration-aware prediction pipeline.
- An explainability protocol covering image attribution and structured clinical features.

7. Comparison of Existing Methods

Method	Image	Clinical	Global Context	Self-Supervised	XAI	External Validation
Conventional CNN	Yes	No	Limited	Usually no	Optional	Variable
ResNet/DenseNet	Yes	No	Moderate	Usually no	Optional	Variable
U-Net/3D U-Net	Yes	No	Moderate	Optional	Optional	Variable
Vision Transformer	Yes	No	Strong	Optional	Possible	Variable

8. Research Gap

The literature indicates several unresolved gaps:

Gap 1: Limited Multimodal Integration:

Many biomedical imaging models remain image-centric despite the availability of clinical information.

Gap 2: Limited Label Efficiency

Supervised models require substantial annotation, while unlabeled medical imaging repositories remain underutilized.

Gap 3: Generalization

High internal test performance does not necessarily imply external clinical generalization.

Gap 4: Missing Modality Robustness

Clinical datasets frequently contain incomplete

- A comprehensive evaluation strategy including external validation, missing-modality testing, robustness, subgroup analysis, and ablation studies.
- A reproducibility specification covering data partitioning, preprocessing, software, hardware, hyperparameters, and statistical analysis.

6. Related Work

CNN-based biomedical image analysis has established strong baselines for classification and segmentation because convolution efficiently captures local patterns. Residual and densely connected architectures improved optimization and feature reuse. Encoder–decoder models such as U-Net became highly influential for biomedical segmentation. Transformer architectures subsequently introduced global token-to-token interactions, while hybrid models combine convolutional locality with attention-based context.

Self-supervised medical representation learning seeks to reduce dependence on labels through masked reconstruction, contrastive learning, or related objectives. Multimodal medical AI combines imaging with clinical variables, reports, or EHR information. Nevertheless, heterogeneous modality scales, missing data, modality dominance, and domain shift remain challenging. Explainability methods such as Grad-CAM, Integrated Gradients, and SHAP can provide evidence, but explanation quality should be quantitatively evaluated.

CNN–Transformer	Yes	Usually no	Strong	Optional	Possible	Variable
Image + clinical concatenation	Yes	Yes	Moderate	Optional	Possible	Variable
Multimodal attention model	Yes	Yes	Strong	Optional	Possible	Variable
MedFusion-X	Yes	Yes	Strong	Yes	Integrated	Core objective

information.

Gap 5: Interpretability

Many models report performance without systematically validating explanation quality.

Gap 6: Calibration

Probability estimates are not always calibrated, which can be problematic for clinical decision support.

Gap 7: Unified Framework

Few frameworks simultaneously address representation learning, multimodal fusion, generalization, explainability, and uncertainty.

The key gap is not the absence of CNNs, Transformers, multimodal learning, or explainability individually. The gap lies in integrating these components into one

experimentally testable framework while simultaneously evaluating label efficiency, cross-domain generalization, missing modalities, calibration, explanation quality, uncertainty, and clinical relevance. MedFusion-X is designed around this integrated evaluation philosophy.

9. Proposed MedFusion-X Framework

The pipeline consists of patient-level curation, modality harmonization, image quality control, preprocessing, self-supervised pretraining, CNN local encoding, Transformer global encoding, clinical feature encoding, cross-modal attention, adaptive fusion, downstream adaptation, uncertainty estimation, calibration, and explainability analysis.

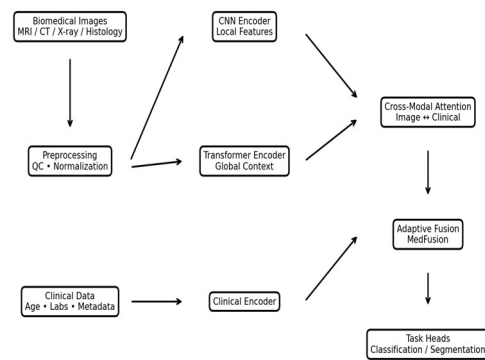
Image representation: $Z_I = F_CNN(I) \square F_TR(I)$, where CNN features capture local spatial patterns and Transformer features capture global contextual dependencies.

Clinical representation: $Z_C = F_C(C)$, where numerical variables are normalized, categorical variables are embedded, and heterogeneous structured information is projected into a common representation space.

Cross-modal attention: $Q_I = Z_I W_Q$, $K_C = Z_C W_K$,

Figure 1. MedFusion Multimodal Biomedical Image Assessment Framework

Figure 1. MedFusion Multimodal Biomedical Image Assessment Framework



Overall architecture linking biomedical images, preprocessing, CNN and Transformer encoders, clinical encoding, cross-modal attention, adaptive fusion and task-specific outputs.

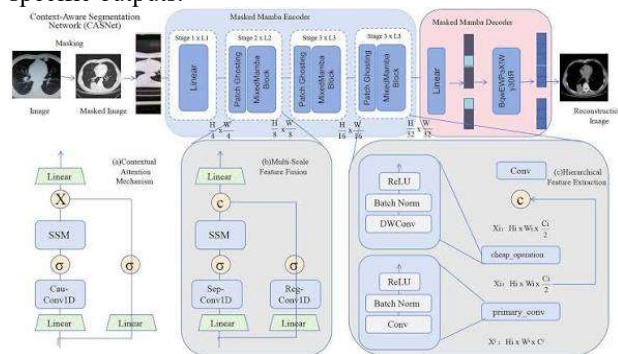


Figure 2. MedFusion Experimental Workflow

$V_C = Z_C W_V$, followed by attention proportional to $\text{softmax}(Q_I K_C^T / \sqrt{d})$.

Adaptive fusion: $Z_F = \alpha_I Z_I + \alpha_C Z_C$, with learned modality weights constrained by a softmax.

Downstream prediction: $\hat{y} = H\psi(Z_F)$, where $H\psi$ is selected for classification, segmentation, detection, grading, or risk prediction.

10. Self-Supervised Representation Learning

The pretraining stage can exploit unlabeled images using masked image modeling, contrastive learning, augmentation consistency, or multimodal alignment. A multimodal objective can additionally encourage alignment between image and clinical representations. Labels or future outcomes must not enter pretraining if the evaluation protocol treats them as unavailable at deployment. The final study should report the exact pretraining objective, augmentation strategy, duration, and data volume.

Figures 1–4 are conceptual architecture illustrations. Figures 5–8 are illustrative placeholders and results generated from actual experiments.

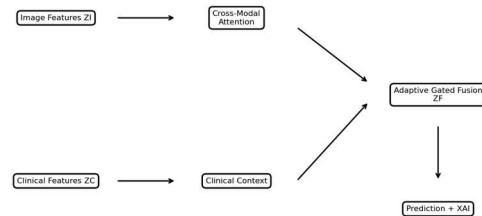
Figure 2. MedFusion Experimental Workflow



End-to-end workflow from patient-level data acquisition and partitioning through preprocessing, multimodal representation learning, prediction, explainability and uncertainty.

Figure 3. Cross-Modal Attention and Adaptive Fusion

Figure 3. Cross-Modal Attention and Adaptive Fusion



Conceptual representation of image-clinical alignment and adaptive multimodal feature fusion.

Figure 4. Explainable AI and Clinical Interpretability Pipeline

Figure 4. Explainable AI and Clinical Interpretability Pipeline

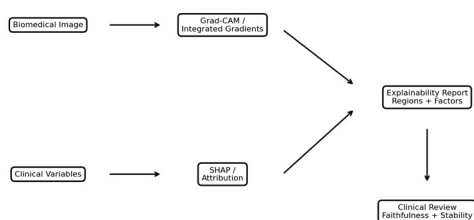


Image attribution and clinical-feature attribution leading to an explainability report and clinical interpretability assessment.

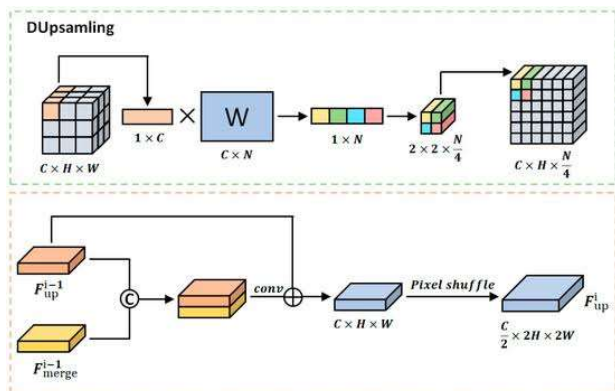
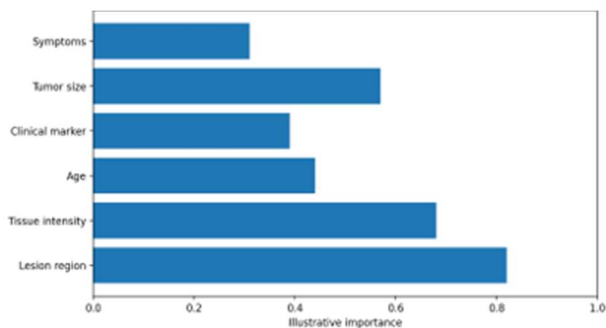


Figure 5. Illustrative Performance Metrics
Synthetic

placeholder visualization only; replace with actual measured performance.

Figure 6. Illustrative Feature Importance Analysis

Figure 6, Illustrative Performance Metrics-
Experimental Results



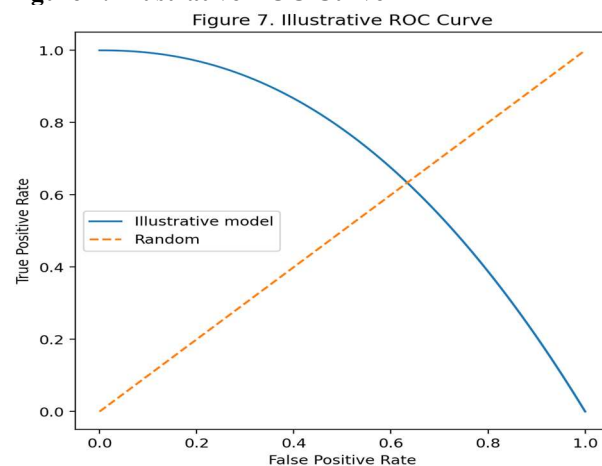
1. Algorithm Objective

Algorithm 1 presents the explainable machine-learning assessment pipeline used within the proposed MedFusion framework for biomedical image analysis. The algorithm integrates biomedical images with structured clinical information, learns local and global image representations, performs multimodal feature fusion, generates a diagnostic prediction, estimates uncertainty,

Importance – Experimental results

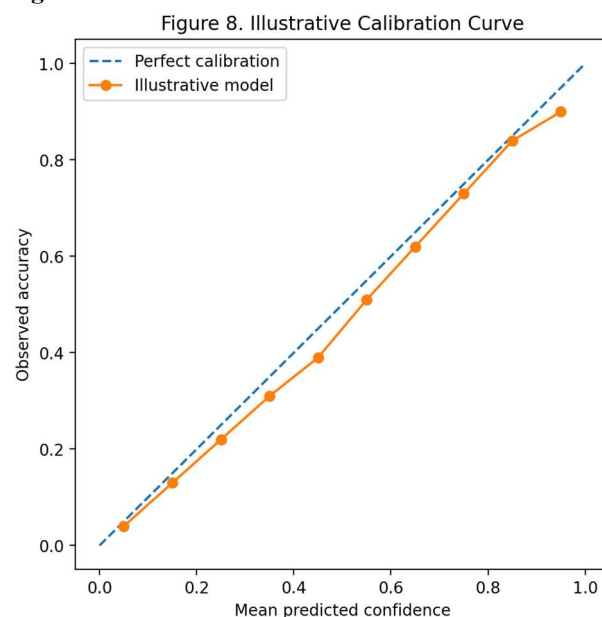
Synthetic placeholder visualization only; replace with actual SHAP or other validated attribution results.

Figure 7. Illustrative ROC Curve



Synthetic placeholder visualization only; replace with ROC data calculated from the real patient-level test cohort.

Figure 8. Illustrative Calibration Curve



Synthetic placeholder visualization only; replace with actual reliability/calibration results.

and produces explainability outputs. The framework can be adapted for MRI, CT, X-ray, ultrasound, histopathology, retinal imaging, and other biomedical image modalities.

The algorithm is intended as a research framework. It does not by itself establish clinical validity. Actual performance values, clinical claims, and statistical conclusions must be generated from appropriately

designed experiments using real datasets.

2. Inputs and Outputs

Category	Definition
Input 1	Biomedical image I: MRI, CT, X-ray, ultrasound, histopathology, etc.
Input 2	Clinical information C: age, sex, laboratory values, symptoms and structured metadata.
Input 3	Optional medical text/report T.
Input 4	Trained MedFusion model M.
Output 1	Predicted disease/class $\hat{Y}$.
Output 2	Prediction probability $P(Y I,C,T)$.
Output 3	Image explanation map EI.
Output 4	Clinical-feature importance EC.
Output 5	Prediction uncertainty U.

3. Algorithm 1 – Data Sets

Algorithm 1: Explainable Machine Learning Biomedical Image Assessment Model

Input: Biomedical image (I), clinical information (C), optional labels (Y)

Output: Prediction ($\hat{Y}$), probability (P), uncertainty (U), explanation (E)

- Step 1:** Group all observations by patient.
- Step 2:** Create patient-level training, validation, and test partitions.
- Step 3:** Perform image quality control and preprocessing.
- Step 4:** Normalize and encode clinical variables.
- Step 5:** Pretrain image representations using a self-supervised learning objective.
- Step 6:** Extract local features using the CNN encoder.
- Step 7:** Extract global contextual representations using the Transformer encoder.
- Step 8:** Encode clinical variables using the clinical encoder.
- Step 9:** Apply cross-modal attention.

Step 10: Compute adaptive multimodal fusion.

Step 11: Generate the fused representation (Z_F).

Step 12: Apply the downstream classification, segmentation, or risk-prediction head.

Step 13: Estimate predictive uncertainty.

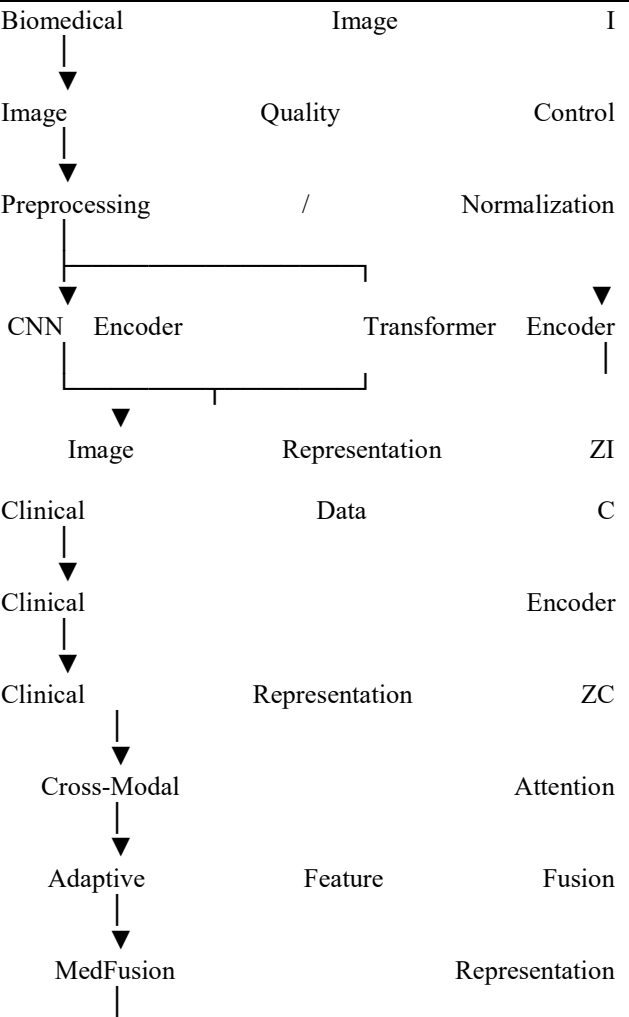
Step 14: Generate image-level explanations using Grad-CAM/Integrated Gradients or an appropriate XAI method.

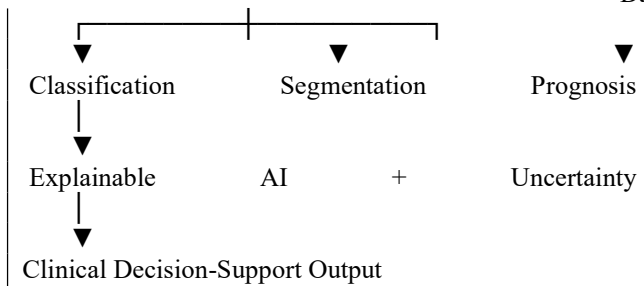
Step 15: Generate clinical feature explanations using SHAP or an appropriate attribution method.

Step 16: Evaluate performance using independent test data.

Step 17: Perform external validation, ablation, calibration, subgroup, and robustness analysis.

4. MedFusion Processing Workflow





5. Mathematical Formulation

Let I represent the biomedical image and C represent structured clinical information. The image encoder produces a latent image representation ZI , while the clinical encoder produces ZC .

$$ZI = FI(I)$$

$$ZC = FC(C)$$

6. Biomedical Image Preprocessing

The preprocessing stage must be adapted to the imaging modality. For MRI/CT, preprocessing may include DICOM/NIfTI loading, orientation normalization, voxel-spacing harmonization, skull stripping where scientifically justified, intensity normalization, registration when required, region-of-interest extraction, and resizing or patch extraction. For X-ray or histopathology, preprocessing may instead involve resolution normalization, artifact handling, color/intensity normalization, and clinically appropriate augmentation. Preprocessing parameters must be estimated using training data only wherever the operation could leak information from the test cohort. The same fixed preprocessing pipeline should then be applied to validation and test data.

7. CNN and Transformer Feature Extraction

The CNN component learns hierarchical local representations such as edges, textures, tissue boundaries and lesion patterns. A Transformer component complements the CNN by modeling long-range dependencies between spatial regions. This hybrid design is intended to capture both fine-grained anatomical characteristics and global structural context.

For volumetric MRI/CT, the implementation can use 3-D CNNs, 3-D Vision Transformers, Swin Transformer variants, slice-based encoders, or a hybrid 2.5-D design depending on computational resources and dataset

Explanation	Primary target	Recommended evaluation
Grad-CAM	CNN feature regions	Localization/faithfulness/stability
Integrated Gradients	Image or	Attribution completeness/stability

10. Uncertainty Estimation

A clinically responsible MedFusion system should distinguish confident predictions from uncertain cases. Predictive entropy, Monte Carlo dropout, deep ensembles, or calibrated probabilities can be used. High uncertainty

The CNN extracts local spatial information and the Transformer models long-range relationships. Their outputs can be combined into a unified image representation:

$$ZI = \text{FusionCNN}(\text{FCNN}(I), \text{FTransformer}(I))$$

Cross-modal attention aligns the imaging and clinical representations:

$$ZA = \text{CrossModalAttention}(ZI, ZC)$$

Adaptive fusion generates the final multimodal representation:

$$ZF = \text{FF}(ZI, ZC, ZA)$$

The task-specific prediction is then:

$$Y_{\text{hat}} = \text{Softmax}(W ZF + b)$$

A composite training objective may be defined as:

$$L_{\text{total}} = \lambda_1 L_{\text{task}} + \lambda_2 L_{\text{align}} + \lambda_3 L_{\text{SSL}} + \lambda_4 L_{\text{re}}$$

characteristics.

8. Multimodal Clinical Fusion

Clinical information can contain complementary evidence that may not be visible in an image alone. Numerical features can be standardized, categorical variables encoded, and missing values explicitly represented. The clinical encoder maps these variables into a latent representation. Cross-modal attention then learns interactions between image and clinical representations rather than assuming that all modalities contribute equally.

Adaptive fusion can assign learned weights to modality representations. This is particularly useful when one modality is noisy, incomplete, or unavailable. Missing-modality experiments should be included in the final MedFusion evaluation.

9. Explainable AI Component

The image explanation module can use Grad-CAM, Grad-CAM++, Integrated Gradients, occlusion sensitivity, or related attribution methods. The resulting heatmap indicates image regions that contribute to the selected model output under the chosen explanation method.

For clinical variables, SHAP, permutation importance, Integrated Gradients, or another suitable attribution technique can estimate the contribution of individual variables. Explanations should be evaluated for stability and faithfulness rather than simply displayed as qualitative heatmaps.

	clinical features	
SHAP	Clinical variables	Feature importance/consistency
Occlusion	Image regions	Prediction change after masking

can be used as a research mechanism for identifying cases that require additional review, but uncertainty thresholds must be validated and should not be treated as a substitute for clinical judgment.

$$H(p) = - \sum_k p_k \log(p_k)$$

11. Performance Evaluation

For classification, report accuracy, precision, recall/sensitivity, specificity, F1-score, ROC-AUC and PR-AUC. PR-AUC and class-wise sensitivity are particularly important for imbalanced biomedical datasets. Calibration should include reliability diagrams, expected calibration error and preferably Brier score. For

12. Python Implementation

```
# MedFusion Explainable Biomedical Image
# Assessment

# Research implementation

import torch
import torch.nn as nn
import torch.nn.functional as F

class MedFusionAssessmentModel(nn.Module):
    def __init__(self, image_encoder, clinical_dim,
                  image_dim, num_classes=2):
        super().__init__()

        self.image_encoder = image_encoder

        self.clinical_encoder = nn.Sequential(
            nn.Linear(clinical_dim, 64),
            nn.ReLU(),
            nn.Dropout(0.2),
            nn.Linear(64, 128),
            nn.ReLU()
        )

        self.image_projection = nn.Linear(
            image_dim, 128
        )

        self.cross_modal_fusion = nn.Sequential(
            nn.Linear(256, 256),
            nn.ReLU(),
            nn.Dropout(0.3),
            nn.Linear(256, 128),
            nn.ReLU()
        )

        self.classifier = nn.Linear(
            128, num_classes
        )

    def forward(self, image, clinical):
        # Extract image representation
        image_features = self.image_encoder(image)

        if image_features.ndim > 2:
            image_features = torch.flatten(
                image_features, 1
            )

        image_features = self.image_projection(
            image_features
```

segmentation, report Dice, IoU/Jaccard, sensitivity and HD95 where appropriate.

Confidence intervals should preferably be calculated using patient-level bootstrap resampling. Statistical comparisons should be paired at the patient level when the same cases are evaluated by competing models.

```
)

# Extract clinical representation
clinical_features = self.clinical_encoder(
    clinical
)

# Multimodal fusion
fused = torch.cat(
    [image_features, clinical_features],
    dim=1
)

fused_features = self.cross_modal_fusion(
    fused
)

logits = self.classifier(
    fused_features
)

probabilities = F.softmax(
    logits, dim=1
)

return logits, probabilities

def predictive_entropy(probabilities):
    """Estimate uncertainty from class probabilities."""
    p = torch.clamp(probabilities, 1e-8, 1.0)
    return -(p * torch.log(p)).sum(dim=1)

def evaluate_prediction(model, image, clinical):
    model.eval()

    with torch.no_grad():
        logits, probabilities = model(
            image, clinical
        )

    prediction = torch.argmax(
        probabilities, dim=1
    )

    uncertainty = predictive_entropy(
        probabilities
    )

    return prediction, probabilities, uncertainty
```

```
#                                     Example:
# prediction, probability, uncertainty = \
#     evaluate_prediction(model, image, clinical)
```

13. Grad-CAM Integration

For a CNN-based image encoder, Grad-CAM should be attached to the final convolutional layer. The gradients of the selected class score with respect to the feature maps are spatially pooled to obtain channel weights. The weighted feature maps are combined and passed through ReLU to produce the explanation map.

```
# Conceptual Grad-CAM calculation
gradients = d(score_class) / d(feature_maps)
weights = mean(gradients, spatial_dimensions)
CAM = ReLU(sum(weights * feature_maps))
CAM = resize(CAM, original_image_size)
```

12. Evaluation Setup and Dataset Strategy

The final study should use datasets appropriate to the biomedical application. For brain MRI, candidate resources include BraTS, ADNI, OASIS, and IXI; for chest imaging, CheXpert and MIMIC-CXR may be considered; ISIC is relevant to dermatological imaging; and TCIA provides multiple oncology imaging collections. Dataset selection must match the research

13. Evaluation Design

Level	Purpose	Evidence of the research
Training	Representation learning	Loss and convergence
Validation	Model selection	Validation metrics
Internal test	Locked performance	Full metrics with CI
External test	Generalization	Independent cohort
Limited-label	Label efficiency	Performance vs label fraction
Missing-modality	Incomplete data	Performance by missingness
Shift robustness	Domain resilience	Performance degradation
Ablation	Component contribution	Matched comparisons
Subgroup	Heterogeneity	Stratified metrics
Calibration	Reliability	ECE/Brier/reliability plot

15. Performance Evaluation

Classification should report accuracy, precision, sensitivity/recall, specificity, F1-score, ROC-AUC, and PR-AUC. For imbalanced datasets, PR-AUC and class-wise sensitivity should receive particular attention. Calibration should be assessed with reliability diagrams,

14. Output Analysis

Output	Description	Used in research
Prediction	Predicted disease/class	Primary diagnostic result
Probability	Class probability	Confidence/calibration analysis
Uncertainty	Predictive entropy or ensemble uncertainty	Reliability assessment
Grad-CAM	Image heatmap	Visual explanation
Clinical importance	SHAP/attributon values	Clinical-feature explanation
Metrics	Accuracy, precision, recall, specificity, F1, AUC	Quantitative evaluation
Error analysis	False positives/false negatives	Clinical risk analysis

question and licensing conditions.

The patient must be the primary unit of partitioning. All images or slices from one patient must remain in one partition. The preferred design is training, validation, locked internal test, and independent external validation. Exact dataset versions, inclusion/exclusion criteria, preprocessing, labels, missing-data rules, and cohort sizes must be documented.

14. Simulation / Experimental Environment

Component	Recommended Specifications
OS	Ubuntu 22.04 LTS or equivalent
Python	3.10/3.11
Deep learning	PyTorch 2.x
Medical imaging	MONAI, NiBabel, SimpleITK
GPU	NVIDIA CUDA-capable GPU; exact model reported
RAM	32 GB+ recommended
Storage	SSD/NVMe
Optimizer	AdamW
Precision	Mixed precision where validated
Randomness	Fixed seeds
Statistics	Patient-level bootstrap and paired tests

expected calibration error, and preferably the Brier score. For segmentation, Dice, Jaccard/IoU, Hausdorff distance or HD95, and sensitivity should be reported. Confidence intervals should preferably be obtained using patient-level bootstrap resampling.

Accuracy = (TP+TN)/(TP+TN+FP+FN); Precision =

$TP/(TP+FP)$; Recall = $TP/(TP+FN)$; Specificity = $TN/(TN+FP)$; F1 = $2(Precision \times Recall)/(Precision + Recall)$.

16. Performance Metrics Figure

Figure 1 should be generated from actual experimental outputs. A point-and-confidence-interval or grouped comparison can show the CNN baseline, Transformer baseline, hybrid model, image-only model, image-plus-clinical concatenation, and full MedFusion-X. ROC and precision–recall curves should be reported when appropriate. No invented numerical values should be inserted.

17. Interpretations and Outcomes

Interpretation should distinguish discrimination, calibration, robustness, and clinical relevance. A model can have high accuracy but poor minority-class sensitivity or poor calibration. Conversely, a small internal performance improvement may be scientifically important if accompanied by better external validation and robustness. Results should therefore be interpreted across

20. Clinical Interpretability Table

Evidence	Method	Interpretation	Validation
Image region	Grad-CAM / IG	Relevant anatomical/pathological region	Expert review or lesion overlap
Clinical variable	SHAP	Feature contribution	Clinical plausibility
Modality contribution	Ablation	Dependence on modality	Performance change
Prediction probability	Calibration	Reliability of confidence	Reliability diagram /ECE
Uncertainty	Entropy/ensemble/MC methods	Low-confidence cases	Expert assessment
Failure case	Counterfactual /ablation	Potential failure mechanism	Case review

multiple dimensions rather than through one headline metric.

18. Feature Importance Analysis

Image explanations may use Grad-CAM, Grad-CAM++, Integrated Gradients, occlusion, or another suitable method. Clinical features may be assessed using SHAP or permutation importance. Multimodal contribution can be studied by modality ablation. Explanation quality should be tested using faithfulness, stability, deletion/insertion, and localization overlap when expert annotations exist. Attention weights alone should not automatically be described as explanations.

19. Feature Importance Figure

Figure 2 should show representative correct and incorrect predictions with image attribution overlays and corresponding clinical-feature importance. Ground-truth lesions or expert annotations should be visually distinguished from model-generated attribution. Quantitative explanation scores should accompany representative examples.

23. Calibration and Uncertainty

Predictive entropy, ensemble variance, Monte Carlo dropout, or other uncertainty methods can be evaluated. Calibration should be reported separately from discrimination. Reliability diagrams should compare predicted probability with observed frequency. A selective-prediction analysis can test whether the model identifies difficult cases by abstaining on the least certain predictions.

24. Statistical Analysis

Patient-level bootstrap resampling is recommended for 95% confidence intervals. Paired statistical tests should be used when predictions are evaluated on the same patients. DeLong's test can compare correlated ROC-AUCs; McNemar's test can compare paired classification outcomes; bootstrap or permutation methods can compare metric differences. Exact p-values, effect sizes, and confidence intervals should be reported. Statistical significance should not be equated with clinical significance.

25. Computational Complexity

Report parameter count, FLOPs or MACs, peak GPU memory, training time, inference latency, model size, and batch size. For three-dimensional imaging, memory requirements can be substantial. Efficient variants may investigate patch-based processing, mixed precision, checkpointing, pruning, quantization, or knowledge distillation.

26. Error Analysis

False positives, false negatives, uncertain predictions, and external-domain failures should be categorized. Where possible, examine lesion size, disease severity, image quality, demographic subgroup, scanner, and missingness. This analysis can reveal shortcut learning and guide future architectural improvements.

30. Conclusion

MedFusion-X provides an extended framework for multimodal biomedical image and clinical-data fusion by combining CNN local representations, Transformer global context, clinical encoding, cross-modal attention, adaptive fusion, self-supervised learning, uncertainty estimation, and explainability. Its central scientific objective is to test whether integrated representations improve discrimination, label efficiency, external generalization, robustness, calibration, and clinical interpretability. The proposed evaluation strategy emphasizes leakage prevention, external validation, ablation, missing-modality testing, statistical confidence intervals, and reproducibility. Claims of state-of-the-art performance or clinical readiness should only follow successful independent experimental validation.

31. Future Work

- Large-scale self-supervised pretraining across multiple biomedical modalities.
- Integration of radiology reports and medical language models.

- Federated multimodal learning across institutions.
- Continual and domain-adaptive learning.
- Missing-modality-aware architectures.
- Foundation-model pretraining across multiple clinical tasks.
- Prospective multi-center validation.
- Human-AI collaboration studies.
- Formal evaluation of XAI faithfulness and clinical utility.
- Efficient deployment using pruning, quantization, distillation, and edge/cloud optimization.

Acknowledgement: The support provided by the Kasireddy Narayanareddy College of Engineering and Research- an UGC Autonomous Institution-Hyderabad and Rajeev Gandhi Memorial College of Engineering & Technology (Autonomous), Nandyal - 518 501, Andhra Pradesh, India is gratefully acknowledged.

32. References

1. Dr. Kommu Naveen & Dr. RMS Parvathi, Scientific Reports-Springer Nature, AI Powered Multi Feature Fusion Framework for Retrieving Images Using Color, Texture and Shape Descriptors. || ISSN 2984-8687 || © November 2025. SCI-Q1 (2025).
2. He K, Zhang X, Ren S, Sun J. Deep Residual Learning for Image Recognition. Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), 2016, 770–778.
3. Huang G, Liu Z, Van Der Maaten L, Weinberger KQ. Densely Connected Convolutional Networks. Proceedings of CVPR, 2017, 4700–4708.
4. Liu Z, Lin Y, Cao Y, et al. Swin Transformer: Hierarchical Vision Transformer using Shifted Windows. arXiv:2103.14030 (2021).
5. Cao H, Wang Y, Chen J, et al. Swin-Unet: Unet-like Pure Transformer for Medical Image Segmentation. arXiv:2105.05537 (2021).
6. Esteva A, Chou K, Yeung S, et al. Deep learning-enabled medical computer vision. npj Digital Medicine 4, 5 (2021).
7. Takahashi S, Sakaguchi Y, Kouno N, et al. Comparison of Vision Transformers and Convolutional Neural Networks in Medical Image Analysis: A Systematic Review. Journal of Medical Systems 48, 84 (2024).
8. Pu Q, Xi Z, Yin S, Zhao Z, Zhao L. Advantages of transformer and its application for medical image segmentation: a survey. BioMedical Engineering Online 23, 14 (2024).
9. Zeng X, Abdullah N, Sumari P. Self-supervised learning framework application

- for medical image analysis: a review and summary. *BioMedical Engineering OnLine* 23, 107 (2024).
10. Madan S, Lentzen M, Brandt J, et al. Transformer models in biomedicine. *BMC Medical Informatics and Decision Making* 24, 214 (2024).
11. Sindhura DN, Pai RM, Bhat SN, et al. A review of deep learning and Generative Adversarial Networks applications in medical image analysis. *Multimedia Systems* 30, 161 (2024).
12. Murshed OAF, Alnaggar B, Jagadale BN, et al. Efficient artificial intelligence approaches for medical image processing in healthcare: comprehensive review, taxonomy, and analysis. *Artificial Intelligence Review* 57, 221 (2024).
13. Vahdati S, Khosravi B, Mahmoudi E, et al. A Guideline for Open-Source Tools to Make Medical Imaging Data Ready for Artificial Intelligence Applications: A Society of Imaging Informatics in Medicine Survey. *Journal of Imaging Informatics in Medicine* 37, 2015–2024 (2024).
14. Litjens G, Kooi T, Bejnordi BE, et al. A survey on deep learning in medical image analysis. *Medical Image Analysis* 42, 60–88 (2017).
15. Zhou Z, Siddiquee MMR, Tajbakhsh N, Liang J. UNet++: A Nested U-Net Architecture for Medical Image Segmentation. *Deep Learning in Medical Image Analysis and Multimodal Learning for Clinical Decision Support*, 2018, 3–11.
16. Isensee F, Jaeger PF, Kohl SAA, Petersen J, Maier-Hein KH. nnU-Net: a self-configuring method for deep learning-based biomedical image segmentation. *Nature Methods* 18, 203–211 (2021).
17. Dosovitskiy A, Beyer L, Kolesnikov A, et al. An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale. *ICLR* (2021).
18. Selvaraju RR, Cogswell M, Das A, et al. Grad-CAM: Visual Explanations from Deep Networks via Gradient-based Localization. *Proceedings of ICCV*, 2017, 618–626.
19. Sundararajan M, Taly A, Yan Q. Axiomatic Attribution for Deep Networks. *Proceedings of ICML*, 2017, 3319–3328.
20. Guo C, Pleiss G, Sun Y, Weinberger KQ. On Calibration of Modern Neural Networks. *Proceedings of ICML*, 2017, 1321–1330.
21. McKinney W, Sieniek M, Godbole V, et al. International evaluation of an AI system for breast cancer screening. *Nature* 577, 89–94 (2020).
22. Kelly CJ, Karthikesalingam A, Suleyman M, Corrado G, King D. Key challenges for delivering clinical impact with artificial intelligence. *BMC Medicine* 17, 195 (2019).
23. Topol EJ. High-performance medicine: the convergence of human and artificial intelligence. *Nature Medicine* 25, 44–56 (2019).
24. Roberts M, Driggs D, Thorpe M, et al. Common pitfalls and recommendations for using machine learning to detect and prognosticate for COVID-19 using chest radiographs and CT scans. *Nature Machine Intelligence* 3, 199–217 (2021).
25. Willemink MJ, Koszek WA, Hardell C, et al. Preparing medical imaging data for machine learning. *Radiology* 295, 4–14 (2020).
26. G. Litjens et al., “A survey on deep learning in medical image analysis,” *Medical Image Analysis*, vol. 42, pp. 60–88, 2017.
27. O. Ronneberger, P. Fischer, and T. Brox, “U-Net: Convolutional networks for biomedical image segmentation,” *MICCAI*, 2015, pp. 234–241.
28. K. He, X. Zhang, S. Ren, and J. Sun, “Deep residual learning for image recognition,” *IEEE CVPR*, 2016, pp. 770–778.
29. G. Huang, Z. Liu, L. Van Der Maaten, and K. Q. Weinberger, “Densely connected convolutional networks,” *IEEE CVPR*, 2017, pp. 4700–4708.
30. A. Dosovitskiy et al., “An image is worth 16×16 words: Transformers for image recognition at scale,” *ICLR*, 2021.
31. A. Hatamizadeh et al., “UNETR: Transformers for 3D medical image segmentation,” *IEEE/CVF WACV*, 2022, pp. 1748–1758.
32. F. Isensee et al., “nnU-Net: A self-configuring method for deep learning-based biomedical image segmentation,” *Nature Methods*, vol. 18, pp. 203–211, 2021.
33. R. R. Selvaraju et al., “Grad-CAM: Visual explanations from deep networks via gradient-based localization,” *IEEE ICCV*, 2017, pp. 618–626.
34. S. M. Lundberg and S.-I. Lee, “A unified approach to interpreting model predictions,” *NeurIPS*, 2017.
35. C. Guo, G. Pleiss, Y. Sun, and K. Q. Weinberger, “On calibration of modern neural networks,” *ICML*, 2017, pp. 1321–1330.

36. S. Azizi et al., “Big self-supervised models advance medical image classification,” IEEE/CVF ICCV, 2021.
37. B. H. Menze et al., “The multimodal brain tumor image segmentation benchmark (BRATS),” IEEE Transactions on Medical Imaging, vol. 34, no. 10, pp. 1993–2024, 2015.
38. M. T. Ribeiro, S. Singh, and C. Guestrin, “‘Why should I trust you?’: Explaining the predictions of any classifier,” ACM SIGKDD, 2016.
39. E. J. Topol, “High-performance medicine: The convergence of human and artificial intelligence,” Nature Medicine, vol. 25, pp. 44–56, 2019.