



WHITE PAPER

Clinical Model Training Methodology

Data sourcing from de-identified sources, training pipeline, specialty-specific fine-tuning, evaluation metrics, and quality assurance processes.

Published:	April 2025
Document ID:	WP-2025-002
Classification:	Public
Author:	ClinixSummary Research Team

This document is provided for informational purposes only. ClinixSummary is a product of GATMEDI. The information contained herein is proprietary and confidential. Reproduction or distribution without prior written consent is prohibited. For the latest version, visit clinixsummary.ai.



1. Executive Summary

The accuracy of any AI clinical documentation system depends entirely on the quality and methodology of its training process. ClinixSummary employs a rigorous, multi-phase training methodology that combines large-scale de-identified clinical data with specialty-specific fine-tuning and continuous clinician-validated improvement.

This white paper details our approach to data sourcing, model training, evaluation, and the continuous improvement cycle that keeps ClinixSummary's outputs aligned with clinical standards across 40+ medical specialties.

2. Data Sourcing & Ethics

2.1 De-identification Standards

All training data undergoes rigorous de-identification compliant with HIPAA Safe Harbor and Expert Determination methods. Protected Health Information (PHI) is removed through a combination of automated NLP-based detection and human review. No raw patient data is ever used in model training.

2.2 Data Sources

Training corpora are sourced from:

- Licensed de-identified clinical datasets from academic medical centres and data cooperatives.
- Publicly available medical literature, clinical guidelines, and documentation standards.
- Synthetic data generated from clinical templates and validated by practicing clinicians.
- Anonymised aggregate patterns from consenting ClinixSummary users (opt-in only, fully de-identified).

2.3 Ethical Review

All data sourcing and training protocols are reviewed by our internal ethics board, which includes practicing clinicians, data privacy specialists, and independent advisors. We maintain a data provenance registry tracking the origin and processing history of all training data.

3. Training Pipeline

3.1 Base Model Pre-training

ClinixSummary's base language models are pre-trained on large medical text corpora including clinical notes, medical textbooks, journal articles, and clinical guidelines. This establishes deep medical language understanding before any speech-specific training begins.

3.2 ASR Fine-tuning

The automatic speech recognition models are fine-tuned on clinical speech recordings with verified transcriptions. Training data spans diverse accents, speaking styles, and clinical environments (quiet offices, busy clinics, operating theatres).



3.3 Specialty-Specific Adaptation

After base training, models undergo specialty-specific fine-tuning for each supported discipline. This involves:

- Specialty vocabulary augmentation (e.g., dental nomenclature, psychiatric rating scales, veterinary anatomy).
- Documentation pattern training on specialty-appropriate note structures.
- Clinical reasoning calibration using specialty-specific clinical scenarios.
- Output format alignment with discipline-standard documentation templates.

4. Evaluation & Quality Metrics

Model performance is evaluated across multiple dimensions:

- Word Error Rate (WER) — Measured against expert-verified transcriptions on held-out test sets.
- Clinical Accuracy Score (CAS) — Proprietary metric measuring correct identification of diagnoses, medications, and clinical findings.
- Section Completeness Index — Measures whether generated notes include all clinically relevant information from the source audio.
- Clinician Acceptance Rate — Percentage of generated notes accepted by clinicians without substantive edits.
- Specialty Terminology Precision — Accuracy of specialty-specific terms, abbreviations, and scoring systems.

Current performance: ClinixSummary achieves a Clinician Acceptance Rate of >92% across core specialties, with continuous improvement through the Kai-zen feedback loop.

5. The Kai-zen Continuous Improvement Loop

ClinixSummary operates on a weekly model update cycle. Clinician feedback (corrections, edits, and ratings) is aggregated, de-identified, and used to continuously refine model outputs. This creates a virtuous cycle where the system improves with every clinical encounter.

The Kai-zen process includes:

- Automated detection of systematic errors or gaps in specialty coverage.
- Prioritised retraining targeting the most impactful improvement areas.
- A/B testing of model updates against baseline before full deployment.
- Clinician advisory panels for each specialty reviewing model outputs quarterly.

6. Model Governance & Versioning

Every model version is tracked in a comprehensive registry with full lineage including training data composition, hyperparameters, evaluation metrics, and deployment dates. Rollback capability ensures that any model update can be reverted within minutes if quality regressions are detected.



7. Conclusion

ClinixSummary's training methodology is designed for the unique demands of clinical AI: accuracy is non-negotiable, specialty nuance matters, and continuous improvement is essential. Our approach — combining rigorous de-identification, specialty-specific fine-tuning, multi-dimensional evaluation, and continuous clinician feedback — ensures that every model update makes the system more accurate, more complete, and more useful for the clinicians who depend on it.