

Model Facts		Model name:		Locale:																																																			
Approval Date:		Last Update:		Version:																																																			
Developer Name:		Developer Contact (phone and email):																																																					
Summary																																																							
Mechanism ▪ Outcome ▪ Output ▪ Target population ▪ Time of prediction ▪ Input data source..... ▪ Input data type ▪ Sensitive attributes included as input data ▪ Training data location and time-period [include information about inclusion / exclusion criteria] ▪ Model type.....																																																							
Validation and performance Example type info <table border="1"> <thead> <tr> <th></th> <th>Prevalence</th> <th>AUC</th> <th>PPV</th> <th>Sensitivity</th> <th>Cohort Type</th> <th>Cohort URL / DOI</th> </tr> </thead> <tbody> <tr> <td>Training Cohort</td> <td></td> <td></td> <td></td> <td></td> <td></td> <td></td> </tr> <tr> <td>Local Retrospective</td> <td></td> <td></td> <td></td> <td></td> <td></td> <td></td> </tr> <tr> <td>Local Temporal</td> <td></td> <td></td> <td></td> <td></td> <td></td> <td></td> </tr> <tr> <td>Local Prospective</td> <td></td> <td></td> <td></td> <td></td> <td></td> <td></td> </tr> <tr> <td>External Cohort</td> <td></td> <td></td> <td></td> <td></td> <td></td> <td></td> </tr> <tr> <td>Target Population</td> <td></td> <td></td> <td></td> <td></td> <td></td> <td></td> </tr> </tbody> </table>								Prevalence	AUC	PPV	Sensitivity	Cohort Type	Cohort URL / DOI	Training Cohort							Local Retrospective							Local Temporal							Local Prospective							External Cohort							Target Population						
	Prevalence	AUC	PPV	Sensitivity	Cohort Type	Cohort URL / DOI																																																	
Training Cohort																																																							
Local Retrospective																																																							
Local Temporal																																																							
Local Prospective																																																							
External Cohort																																																							
Target Population																																																							
Risk mitigations to ensure fairness ▪ Problem identification and procurement steps: ▪ Development and adaptation steps: ▪ Clinical integration steps: ▪ Lifecycle management steps:																																																							
Uses and directions ▪ Intervention prompted by model output: ▪ Benefits: ▪ Target population and use case: ▪ Intended user of model: ▪ General use: [provide information about whether the model informs, augments, or replaces clinical judgment] ▪ Appropriate decision support: ▪ Before using this model: ▪ Safety and efficacy evaluation:																																																							
Warnings ▪ Risks: ▪ Inappropriate Settings: ▪ Clinical Rationale: ▪ Inappropriate decision support: ▪ Generalizability: ▪ Discontinue use if:																																																							
Post-implementation performance monitoring ▪ Model and intervention monitoring: ▪ Intervention effectiveness monitoring:																																																							

- **Model fairness monitoring:**
- **Intervention fairness monitoring:**
- **Intervention updating criteria:**
- **Model updating criteria:**

Other information:

- **Outcome Definition:**
- **Related model:**
- **Intervention treatment effect estimation:** [provide references and links to publications that demonstrate treatment effect of intervention prompted by model]
- **Model development & validation:**
- **Model implementation:**
- **External validation methods and results:**
- **Clinical trial:**
- **Clinical impact evaluation:**
- **Funding source for model development:**
- **Party that conducted external testing:**

Open Access License

This work builds off prior work performed by Sendak, et. al related to the 'Model Facts' label and published in Nature Digital Medicine:

<https://www.nature.com/articles/s41746-020-0253-3>. The prior work was released under a Creative Commons Attribution 4.0 International License:

<https://creativecommons.org/licenses/by/4.0/>. Changes were made to the 'Model Facts' label to include all data elements required for HTI-1 compliance.

This 'Model Facts' label is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license, and indicate if changes were made. To view a copy of this license, visit <http://creativecommons.org/licenses/by/4.0/>.

MODEL FACTS | FICTIONAL EXAMPLE

Model name: ChartSage Review Assist | Locale: U.S. ambulatory community health centers

Approval Date: May 15, 2026 (fictional)

Last Update: August 1, 2026 (fictional)

Version: 1.2-demo

Developer Name: ClearView Clinical AI Lab (fictional)

Developer Contact: modelgovernance@example.invalid | 555-010-2040 (fictional; nonworking)

Summary

- ChartSage Review Assist is a fictional clinician-facing AI chart review tool that organizes longitudinal EHR data and highlights potential preventive, chronic-care, medication-reconciliation, and follow-up opportunities for human review. It does not diagnose, order, prescribe, message patients, or modify the health record autonomously.
- **Synthetic-data notice:** Every organization, date, contact, cohort, metric, and result in this completed card is invented for demonstration and must not be used as evidence of real-world safety or performance.

Mechanism

- **Outcome:** Presence of at least one predefined chart-review opportunity confirmed by two trained clinical reviewers using an adjudication guide. Categories include preventive care, chronic disease monitoring, medication reconciliation, abnormal-result follow-up, and care-plan documentation.
- **Output:** A 0-100 review-priority score, category-specific flags, confidence band, and links to the source encounters, laboratory results, medications, and documentation used to generate each flag. No patient-facing output.
- **Target population:** Patients age 18 years and older with at least one ambulatory encounter in the prior 24 months. Excludes test patients, records marked deceased before the review date, patients who opted out of secondary model use, and charts lacking minimum encounter history.
- **Time of prediction:** On demand before a scheduled visit or during an authorized population-health chart review. Data cutoff is displayed to the user.
- **Input data source:** Structured EHR tables and selected clinical-note sections from an authorized longitudinal ambulatory record.
- **Input data type:** Demographics; encounter history; diagnoses; procedures; medications; allergies; immunizations; laboratory results; vital signs; referrals; care gaps; and bounded note text from assessment/plan, medication reconciliation, and result follow-up sections.
- **Sensitive attributes included as input data:** Age is used for guideline applicability. Sex assigned at birth is used only where a reviewed rule requires it. Race, ethnicity, preferred language, disability status, gender identity, sexual orientation, ZIP code, and payer are excluded from scoring but retained in a separate protected audit dataset for fairness evaluation.
- **Training data location and time-period:** Synthetic multi-site U.S. ambulatory dataset representing 2019-2024 care patterns. Inclusion: adults meeting the target-population criteria with sufficient longitudinal history. Exclusion: test/deceased/opt-out records, unsupported document types, and charts with unresolved identity matching. All details are fictitious.
- **Model type:** Hybrid retrieval and classification system: deterministic eligibility rules, clinical concept extraction, gradient-boosted category classifiers, and a constrained summarization component that may only cite retrieved chart evidence. Outputs are calibrated and thresholded by category.

Validation and performance - FICTIONAL ILLUSTRATIVE RESULTS

Cohort	Prevalence	AUC	PPV	Sensitivity	Cohort Type	Cohort URL / DOI
Training Cohort	24.0%	0.89	0.63	0.82	Synthetic development; n=120,000 reviews	N/A - synthetic dataset
Local Retrospective	22.7%	0.86	0.58	0.79	Held-out retrospective; n=18,500	N/A - fictional internal evaluation
Local Temporal	21.9%	0.84	0.55	0.76	Later-period holdout; n=9,200	N/A - fictional internal evaluation
Local Prospective	23.5%	0.82	0.52	0.74	Silent prospective; n=4,800	N/A - fictional internal evaluation

External Cohort	26.1%	0.80	0.56	0.71	Synthetic external-site holdout; n=11,300	N/A - synthetic dataset
Target Population	23.1%	0.83	0.55	0.75	Pooled intended-use evaluation; n=25,300	N/A - illustrative pooled result

Metrics below are synthetic examples. Thresholds were selected separately for each cohort before evaluation; values are not evidence of clinical effectiveness.

Risk mitigations to ensure fairness

- **Problem identification and procurement steps:** Define the clinical problem and non-AI alternatives; document intended users and affected populations; review accessibility, privacy, data rights, security, and subgroup evidence; require disclosure of limitations, known failure modes, update policy, and audit access; and reject use cases that automate eligibility, coverage, or disciplinary decisions.
- **Development and adaptation steps:** Exclude social and protected attributes from scoring unless clinically necessary and approved; assess intersectional performance across age, race, ethnicity, preferred language, sex assigned at birth, disability indicator, payer, rurality, and data-completeness strata; calibrate thresholds by clinical harm rather than equalizing a single metric; conduct clinician and patient-advisor review of flag language.
- **Clinical integration steps:** Display source evidence, confidence, and data recency; require clinician acceptance or dismissal with optional reason; suppress unsupported flags; use accessible language and keyboard navigation; prohibit scores from being copied as diagnoses; provide escalation for suspected systematic bias.
- **Lifecycle management steps:** Review performance and fairness monthly during pilot and quarterly thereafter; use minimum sample-size and confidence-interval rules; investigate meaningful subgroup gaps; maintain change control, model/data lineage, rollback, incident response, and retirement criteria; publish approved updates to users.

Uses and directions

- **Intervention prompted by model output:** A trained reviewer verifies the cited evidence, checks current guideline and patient context, then may add an item to the pre-visit planning workflow, reconcile information, route a question to the responsible clinician, or dismiss the flag.
- **Benefits:** May reduce manual search burden, improve consistency of chart review, and make overlooked documentation or follow-up needs easier to locate. These are intended benefits, not demonstrated real-world claims.
- **Target population and use case:** Authorized adult ambulatory chart review for pre-visit planning, quality-improvement review, and clinically supervised population-health workflows.
- **Intended user of model:** Licensed clinicians, nurses, pharmacists, medical assistants, and trained quality or population-health staff acting within role-based permissions and local policy.
- **General use:** Augments clinical judgment. It does not replace independent review, patient conversation, guideline interpretation, or clinician accountability.
- **Appropriate decision support:** Prioritize which charts or chart sections to review; surface potentially relevant evidence; support a documented human review decision.
- **Before using this model:** Confirm patient identity, review timestamp and data cutoff, validate that the encounter and workflow are in scope, inspect linked evidence, consider outside records and patient preferences, and follow local escalation policy.
- **Safety and efficacy evaluation:** Before activation, complete privacy/security review, workflow hazard analysis, usability testing, retrospective validation, subgroup evaluation, and a silent prospective pilot. Do not infer clinical benefit until evaluated in a governed impact study.

Warnings

- **Risks:** False positives may create workload or unnecessary follow-up; false negatives may miss a review opportunity; incomplete or stale records may mislead; extracted note context may be wrong; automation bias may lead users to over-trust a flag; subgroup performance may differ.
- **Inappropriate settings:** Emergency triage, real-time diagnosis, inpatient deterioration detection, autonomous ordering or prescribing, patient denial or eligibility decisions, billing enforcement, employee evaluation, research recruitment without approval, and any use outside authorized ambulatory workflows.
- **Clinical rationale:** Output indicates review priority, not disease probability, care quality, negligence, or a definitive care gap. A displayed flag must be interpreted with the linked source data and current clinical context.

- **Inappropriate decision support:** Do not use as the sole basis for diagnosis, treatment, medication change, outreach, referral closure, patient risk labeling, performance management, reimbursement, or resource denial.
- **Generalizability:** Illustrative results may not transfer to pediatric, inpatient, specialty, non-U.S., low-data, or substantially different EHR environments. Local validation is required after meaningful workflow, coding, or population changes.
- **Discontinue use if:** Linked evidence cannot be retrieved; data mapping changes without validation; critical subgroup degradation or unsafe failure patterns are detected; calibration exceeds approved limits; unauthorized output exposure occurs; or rollback is directed by governance, privacy, security, or clinical safety leadership.

Post-implementation performance monitoring

- **Model and intervention monitoring:** Track data completeness, extraction failures, score distribution, alert volume, acceptance/dismissal, override reasons, latency, and suspected incidents by site and workflow. Review monthly during pilot and quarterly after approval.
- **Intervention effectiveness monitoring:** Assess reviewer time, verified-opportunity yield, downstream completion of clinically appropriate actions, duplicate work, patient complaints, and balancing measures such as unnecessary tests or outreach. Use comparison designs where feasible.
- **Model fairness monitoring:** Evaluate sensitivity, PPV, false-negative rate, calibration, and missing-data rates across approved demographic and access-related audit groups, including intersectional groups when sample size supports reliable interpretation.
- **Intervention fairness monitoring:** Evaluate who receives review, outreach, completed follow-up, and potential burdens after human decisions. Examine differences in dismissal and override reasons rather than attributing disparities to the model alone.
- **Intervention updating criteria:** Update workflow instructions or thresholds when burden, benefit, safety, accessibility, or equity measures are outside governance-approved limits, or when clinical guidance or operational ownership changes.
- **Model updating criteria:** Retrain or recalibrate after material coding, terminology, EHR-mapping, population, prevalence, or documentation shifts; sustained performance degradation; new evidence; or a significant data-quality incident. Require versioned validation and approval before release.

Other information

- **Outcome Definition:** A chart-review opportunity is positive only when adjudicators identify sufficient chart evidence that a predefined review step is warranted at the index date. Disagreements are resolved by a senior clinical reviewer. The label is not a diagnosis or a measure of clinician performance.
- **Related model:** None. A separate deterministic eligibility rules engine may supply guideline applicability but does not share the model score.
- **Intervention treatment effect estimation:** Not established. No real publication or clinical trial is claimed for this fictional tool. A stepped-wedge or cluster-randomized evaluation could estimate benefits and balancing harms before scaled deployment.
- **Model development & validation:** Fictional synthetic-data development with patient-level splits, site holdouts, a later-time holdout, locked thresholds, calibration assessment, subgroup analyses, error review, and clinician usability testing.
- **Model implementation:** Role-based access; on-demand or batch pre-visit workflow; evidence provenance; human sign-off; audit logging; feedback capture; version display; rollback; and local governance approval.
- **External validation methods and results:** Illustrative synthetic external-site holdout shown above. No real-world external validation has been performed.
- **Clinical trial:** None. Any prospective study would require appropriate ethics, privacy, safety, and operational approvals.
- **Clinical impact evaluation:** Planned fictional pilot endpoints: verified review opportunities per 100 charts, median review time, downstream completion, unnecessary-action balancing measures, user-reported workload, and subgroup differences.
- **Funding source for model development:** Internal innovation demonstration budget of ClearView Clinical AI Lab (fictional).
- **Party that conducted external testing:** Independent Synthetic Health Analytics Collaborative (fictional).

DEMONSTRATION ONLY: This completed model card contains fictitious information and synthetic performance data. Replace and validate every field before any operational use.

Open Access License

This work builds off prior work performed by Sendak, et. al related to the 'Model Facts' label and published in Nature Digital Medicine: <https://www.nature.com/articles/s41746-020-0253-3>. The prior work was released under a Creative Commons Attribution 4.0 International License: <https://creativecommons.org/licenses/by/4.0/>. Changes were made to the 'Model Facts' label to include all data elements required for HTI-1 compliance.

This 'Model Facts' label is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license, and indicate if changes were made. To view a copy of this license, visit <http://creativecommons.org/licenses/by/4.0/>.

AI Vendor Questionnaire

Please complete the following questions, providing as much detail as possible.

VENDOR CONTACT INFORMATION

Vendor Name	
Contact Name	
Title	
Address	
Email	
Website	
Summary of Products & Services	

AI Functionality & Features				
Yes	No	N/A	Question	Details (if applicable)
			Does your AI tool have key features and capabilities?	
			Does your AI tool address specific clinical or operational problems?	
			Is your AI's core technology based on deep learning, reinforcement learning, supervised learning, or another specific architecture? Please specify.	
			Can your AI tool handle large data volumes and real-time settings effectively?	
			Will the AI tool's intelligence and accuracy improve over time?	
			Can the AI tool work on data it has not been trained on before?	

			Are specific metrics used to evaluate the effectiveness of your AI tool?	
Patient Safety & Care Quality				
Yes	No	N/A	Question	Details (if applicable)
			Have you consistently demonstrated effective bias management and accuracy in your AI models?	
			Do you have a process to ensure your AI model's accuracy?	
			Do you ensure patient consent and privacy when using the AI tool?	
			Has your AI tool achieved measurable results?	
			Has the accuracy of your AI tool been evaluated?	
			Do you test for bias and accuracy in your AI model?	
			Can you provide evidence or documentation of bias testing and results?	
			Do you have a quality assurance process in place?	
			Are there defined quality control procedures for your AI tool?	
Data Security & Compliance				
Yes	No	N/A	Question	Details (if applicable)
			Is your company facing any pending AI-related litigation, threats of litigation, or regulatory inquiries?	
			Does your AI model maintain the sensitivity and confidentiality of data?	

			Have you obtained the necessary certifications and/or approvals for the AI tool?	
			Does the AI tool comply with relevant health care regulations (e.g., HIPAA, state privacy laws)?	
			Is your company SOC Type 2 certified?	
			Do you have an acceptable use policy?	
			Is data encrypted and stored securely?	
			Does a specific entity or individual have ownership and control over any of the data?	
			Is the data stored in a cloud-based infrastructure, on dedicated on-premises servers or in a hybrid system?	
			Do you have measures in place to ensure data security and privacy?	
			Does your AI solution collect specific types of data?	
			Is there a defined retention period before data is discarded?	
			Do you have a process for handling data breaches and security incidents?	
Integration & Scalability				
Yes	No	N/A	Question	Details (if applicable)
			Does your AI tool integrate with existing systems?	

			Is there a defined process for integrating your AI tool into existing systems?	
			Can a pilot or proof-of-concept project be conducted?	
			Is onboarding effort minimal for users?	
			Does your software/platform have any limitations around number of users, volume of requests, or topics?	
			Do you offer localization services?	
			Do you have a roadmap for future development?	
			Are you shifting from general purpose AI to specialized AI?	
Vendor Expertise & Support				
Yes	No	N/A	Question	Details (if applicable)
			Does your company have a history and expertise in developing AI tools specifically for health care?	
			Can you provide references or client testimonials from health care organizations?	
			Have you identified competitors?	
			Are you building any of your own infrastructure, and does it include specific capabilities?	
			Do you have a process for handling, collecting, and responding to feedback?	

			Is there a structured process for submitting support requests post-implementation?	
			Are ongoing support and maintenance services included?	
			Do you monitor AI tool performance post-deployment?	
			Is there a process for handling AI tool incidents or malfunctions?	
			Do you provide training for users?	
Ethical Considerations & Accountability				
Yes	No	N/A	Question	Details (if applicable)
			Are established ethical guidelines followed in the development and deployment of AI tools?	
			Is there a policy on accountability for errors or negative outcomes from the AI tool?	
			Do you have a public AI ethics statement or policy? If yes, please provide the link.	
			Do you provide liability coverage for AI tool failures or inaccuracies?	
			Do your generative AI models include citations to verify referenced information?	
			Is the system auditable, and can the audit process be customized?	
			Will users have visibility into how the AI operates?	

			Do humans review your AI model's results?	
Training Data & Customization				
Yes	No	N/A	Question	Details (if applicable)
			Is the training data representative of the patient population?	
			Can you confirm that your AI models were trained with properly licensed content?	
			Was your AI model trained using a specific dataset and by identifiable trainers?	
			Were the training methods chosen based on a defined rationale?	
			Can your AI model be customized to our specific needs?	
			Is your AI model updated and retrained over time?	
Cost & ROI				
Yes	No	N/A	Question	Details (if applicable)
			Is the total cost of ownership clearly outlined, including licensing, implementation, and maintenance?	
			Are there ongoing costs such as support, updates, or subscription fees?	
			Are variable costs or rate lock-in options available?	
			Are there additional costs for optimization, future updates, or add-ons?	
			Can the expected ROI of the AI software/platform be quantified?	

AI Vendor Assessment

Use the completed vendor questionnaire to evaluate the vendor's qualifications.

Note: A score of **0** in response to any compliance, regulatory, or security-related questions may warrant disqualification.

2 - Meets requirements: The answer is complete, well-documented, specific, and aligns with the organization's standards.

1 - Needs clarification: The answer is partially complete or vague, suggesting further follow-up is needed. **0 - Does not meet requirements:** The answer is incomplete, non-specific, or fails to sufficiently address the question.

Category	High Priority Questions	0	1	2	Notes
AI Functionality & Features	What are the key features and capabilities of your AI tool?	☐	☐	☐	
Patient Safety & Care Quality	How would you describe your history with bias management and ensuring accuracy in your AI models?	☐	☐	☐	
	How do you ensure your AI model's accuracy? (Request case studies or test results.)	☐	☐	☐	
	How do you ensure patient consent and privacy when using the AI tool?	☐	☐	☐	
Data Security & Compliance	Does your company face any pending AI-related litigation, threats of litigation, or regulatory inquiries?	☐	☐	☐	
	How do you handle sensitive or confidential information?	☐	☐	☐	
Integration & Scalability	How does your AI tool integrate with our existing systems?	☐	☐	☐	
Ethical Considerations & Accountability	What ethical guidelines and principles do you follow in the development and deployment of your AI tools?	☐	☐	☐	
	What is your policy on accountability for errors or negative outcomes resulting from the AI tool?	☐	☐	☐	
Training Data and Customization	How do you ensure that the training data is representative of my patient population?	☐	☐	☐	
	Can you confirm that your AI models were trained with properly licensed content?	☐	☐	☐	
Cost	What is the total cost of ownership of your AI system, including licensing, implementation, and ongoing maintenance?	☐	☐	☐	
	Beyond the initial cost, what are the ongoing costs such as support, updates, or subscription fees?	☐	☐	☐	
SCORING TOTALS					TOTAL

0-8	The vendor does not meet the organization's standards.
9-16	The vendor may meet the standards, but further follow-up and documentation may be necessary.
17-26	The vendor meets the organization's standards.