

Skin Disease Prediction and Medicine Recommendation Using Hybrid CNN-RBM with Doctor Verification

Pokala Bhargavi*, Pilla Srikar, Narem M N S C Ganesh, Althi Vinodh Kumar, Chilla Mokshagna

Dept. of CSE (AI & ML), Anil Neerukonda Institute of Technology & Sciences (ANITS),
Visakhapatnam, India

DOI: <https://doi.org/10.51583/IJLTEMAS.2026.150700035>

Received: 23 July 2026; Accepted: 28 July 2026; Published: 07 August 2026

ABSTRACT

Most automated tools for skin diagnosis look at a photograph and stop there. Symptom data — how long the rash has been present, whether it itches, where on the body it appears — is routinely discarded, even though no practising dermatologist would ignore it. On top of that, AI-generated outputs typically reach patients without any medical review at all. This paper describes a system built to correct both problems. Patients upload a lesion photograph and fill out a twenty-item symptom questionnaire. The image is processed by an EfficientNet-B3 convolutional network [16], which produces a 512-dimensional feature vector. A Restricted Boltzmann Machine condenses the symptom responses into a 100-dimensional latent vector. The two feature vectors are fused into a 612-dimensional representation and passed to a softmax classifier. A dermatologist, after consulting the Doctor Portal, is the only one able to share the output with the patient — this step cannot be skipped.

The platform was built using React for the front-end, Flask for the server, and MongoDB for data storage. The system authenticates and encrypts data using JWT tokens, bcrypt, and AES-256.

Our CNN+RBM model outperforms four baseline models — CNN-only, Matrix Factorisation (MF), SVD, and Weighted SVD (WSVD) — across every metric measured. By epoch 14, the model achieves 72.58% accuracy, 74.1% precision, 71.8% recall, and a 72.9% F1 score, with a false positive rate of only 4.2%.

Keywords: skin disease classification, multimodal deep learning, convolutional neural networks, Restricted Boltzmann Machine, tele-dermatology, doctor verification, DermNet, medicine recommendation

INTRODUCTION

In reality, access to dermatological care is still extremely limited in rural India. The majority of districts have less than one dermatologist per 100,000 people, according to government statistics and WHO estimations, and this has not altered over the last ten years [1]. Because of this, a person's actual alternatives when they find a rare skin ailment are limited: they can wait weeks for a specialist consultation, try an over-the-counter treatment and hope it works, or visit a general practitioner who lacks dermatology competence. Many ailments that may have been easily and affordably addressed therefore deteriorate over time.

Melanoma is the most cited example — it responds well to treatment when caught early and very poorly when caught late [1].

Mobile phones are a different story. Smartphone ownership expanded into the same rural areas that have no dermatologist, often years ahead of any healthcare infrastructure. Several apps have attempted to use this gap. The typical design is simple: take a photo, get a result. What most of these applications miss is that a photograph alone does not tell you what a patient is experiencing. Duration of the condition, level of itching, whether it is spreading, which part of the body is affected — these clinical details matter, and most apps simply do not collect them [2]. Even worse, several tools bypass doctors entirely and send predictions straight to the patient [3]. A person who reads “Psoriasis — 84% confidence” and buys a corticosteroid cream for what is actually a fungal infection has been actively harmed.

Our system was designed with these two failures in mind. A patient provides both a photograph and answers to twenty symptom questions. The two inputs are processed by different models whose outputs are merged before any classification is made. The AI result never goes directly to the patient — it goes to a licensed dermatologist first. The contributions of this work are:

- 1) A picture of the afflicted skin area and the patient's answers to twenty questions about their symptoms are the two forms of input that the system takes.
- 2) The model combines a CNN and an RBM to assess both types of data. Features from the image are combined with specifics of the symptoms to make a diagnosis.
- 3) A doctor's review is required in the Doctor Portal as a mediator between the patient and the AI.
- 4) The system is completely deployed with React, Flask, and MongoDB, using JWT, bcrypt, and AES-256 to keep it secured.

Relevant work is discussed in Section II. The system as a whole is explained in Section III. The AI models are covered in Section IV. Security is the main topic of Section V. The findings are displayed in Section VI. Limitations and the paper's conclusion are covered in Sections VII and VIII.

RELATED WORK

A. Image-Based Skin Disease Classification

The paper most researchers cite as the starting point for CNN-based dermatology is Esteva et al. [4], who trained a deep network on over 129,000 clinical images and reached dermatologist-level performance on two cancer types. Important work, but narrowly scoped and entirely image-based. The HAM10000 dataset [5] later provided a benchmark that allowed fairer comparisons between architectures. Hameed et al. [6] used it to show that residual and dense networks consistently outperform simpler CNN designs. The seminal AlexNet work [20] demonstrated that deep CNNs trained on large image datasets can achieve human-competitive accuracy, a foundation that underpins all subsequent dermatology CNN work. Despite this progress, purely image-based approaches have not closed the gap with clinical diagnostic accuracy across the full range of skin conditions. Kassem et al. [7] reviewed this body of work and concluded that integrating patient-reported symptoms remains the most underexplored direction.

B. Multimodal and Explainable Approaches

Although the results were superior to image-only approaches, scaling was challenging due to the need for coupled image-text data. Chanda et al.'s [9] finding that incorporating Grad-CAM heatmaps into CNN predictions significantly increased the chance that physicians would follow AI suggestions had a significant impact on our Doctor Portal design. Although initially intended for segmentation, the U-Net architecture [19] showed the usefulness of encoder-decoder paths for maintaining spatial detail in medical images. In their assessment of deep learning for healthcare in general, Miotto et al. [18] highlighted the significance of clinical integration and interpretability in implemented systems.

C. Recommendation Systems and Hybrid Architectures

Chinnasamy et al.'s earlier research is the most immediately applicable to our symptom model [10]. They outperformed MF, SVD, and WSVD on both RMSE and MAE when they paired an RBM with CNN-based collaborative filtering for a general healthcare recommendation job. We added the doctor verification layer after applying the same hybrid structure specifically to the prediction of skin diseases. Separately, Khan et al. [11] showed that beginning an RBM with weights pre-trained on structured data yields better results than random initialization. We directly use Hinton et al.'s contrastive divergence training approach, which provided the fundamental theory for RBMs [17]. Aujla et al. [12] showed that it is possible to implement comparable deep

learning recommendation systems in actual healthcare settings. EfficientNet [16] served as the backbone for our CNN, selected for its optimal combination of precision and inference speed on conventional hardware.

D. Clinical Safety and Regulatory Context

A system that scores well on a test benchmark isn't necessarily safe or useful in a real clinic. Topol [13] emphasises the difference: AI in medicine should augment doctors, not replace them. That is the premise of our Doctor Portal. Notably, Obermeyer and Emanuel [14] point out that machine learning can amplify demographic bias, especially for those already excluded from medical datasets. Rajpurkar et al. [15] suggest that good performance on tests is not necessarily what helps doctors with real-world cases in real hospitals. Portugal et al. [3] examined recommendation systems in machine learning and identified wide gaps around transparency and clinical validation. We took all these findings into account and made one major decision: every AI recommendation must be verified by a doctor. That is not merely a recommendation; it is baked into the system.

System Architecture

Four modules make up the system: Patient, Doctor, AI Inference, and Case Management. They communicate via a Flask REST API backed by a single MongoDB database. The important endpoints are described in Table I. Fig. 1 depicts data flow from patient submission all the way through AI processing and doctor evaluation to the end of prescription delivery.

TABLE I. REST API Endpoint Summary

Module	Endpoint	Method	Function
Patient	/register	POST	Create account
Patient	/login	POST	Issue JWT token
Patient	/upload-case	POST	Submit image + symptoms
Patient	/case-status/{id}	GET	Poll case status
Doctor	/pending-cases	GET	Unreviewed case queue
Doctor	/submit-diagnosis	POST	Save diagnosis + Rx
AI (int.)	/predict	POST	Run CNN+RBM inference

Fig. 1. System Architecture — Data Flow Diagram

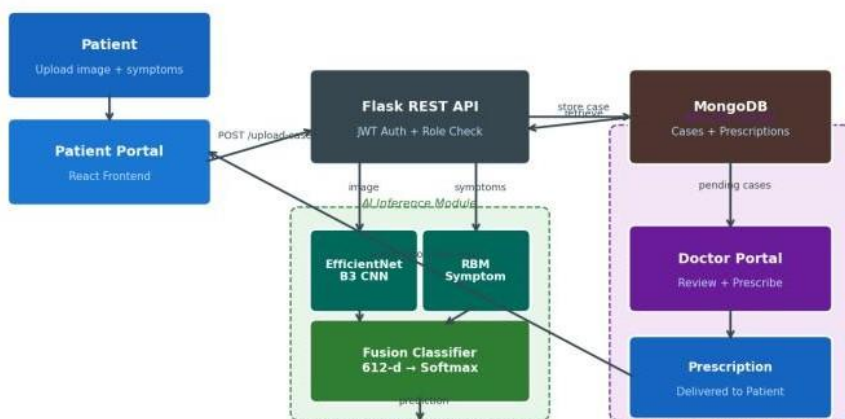


Fig. 1. System Architecture — End-to-End Data Flow Diagram

A. Patient Module

After login, an image upload field and a symptom form with twenty questions about intensity of itch, redness, swelling or discharge, duration, body area, past diagnoses, and current medications are shown to the user. The symptom responses are POSTed to /upload-case as a JSON array, while the photo is sent as multipart form data in the same call. The case is then pending. Pending, Reviewed, and Completed are the possible return values of a poll to GET /case-status/{case_id}. Only after Completed does the prescription appear.

B. Doctor Module

A role-scoped JWT is generated when doctors log in to a separate endpoint. The middleware intercepts a patient token before any application code runs, so it cannot be used for routing to doctor endpoints. The case queue — where each uploaded image, the symptom breakdown, and the AI prediction and confidence score are shown — is the main physician interface. The physician completes the structured POST /submit-diagnosis prescription and may accept or reject the AI recommendation.

C. AI Inference Module

The /predict endpoint is not externally callable; it is triggered internally from /upload-case on every inbound case. The models are loaded on Flask startup, so the first inference call is no slower than subsequent calls. The /predict endpoint receives an image path and symptom vector, runs both models, concatenates the two feature vectors, classifies the combined vector, and outputs a ranking of disease predictions with probabilities.

D. Case Management Module

Each MongoDB case document contains: case_id, user_id, path to the photo image, a twenty-valued float array for symptom values, AI predictions as disease-confidence pairs, doctor_id (empty until reviewed), a status string, and a creation timestamp. Prescriptions are stored in a separate collection and linked to case documents via case_id. A third reference collection maps disease labels to standard treatments used to pre-populate the prescription form for routine cases.

AI Implementation

A. Datasets and Preprocessing

Images were sourced from DermNet, a public dermatology image database spanning a large number of conditions and a range of skin types and body sites. Real patient symptom data was unavailable; this limitation is discussed in Section VII. Instead, 15,000 synthetic symptom records were generated by adding Gaussian noise to the published average clinical profile for each disease to model typical patient variation. All images were resized to 224×224 and normalised to [0, 1]. Augmentation was applied on-the-fly each epoch: horizontal and vertical flips, rotation up to $\pm 25^\circ$, brightness and contrast variation of $\pm 20\%$, and random crop-resize, ensuring the model rarely saw the same image form twice across epochs.

B. CNN — EfficientNet-B3

The image feature extractor is EfficientNet-B3 [16], pre-trained with ImageNet weights [20]. B3 was selected over larger variants (B4–B7) because experiments indicated that the accuracy gain from larger models was not worth the added inference time. The original classification head was replaced by global average pooling and a dense layer with 512 units (ReLU, dropout 0.4), giving the image feature vector $f_e \in \mathbb{R}^{512}$.

Training was carried out in two stages. Phase one (epochs 1–10) froze the EfficientNet backbone and updated only the new dense head, so that noisy gradients would not overwhelm the useful pretrained features early on. In phase two (epochs 11–25), the full network was unlocked for joint fine-tuning at a learning rate ten times lower. Adam with cosine annealing was used throughout, with categorical cross-entropy as the loss function. The best-performing checkpoint on the validation set was saved for inference.

Fig. 7. CNN Training and Validation Loss

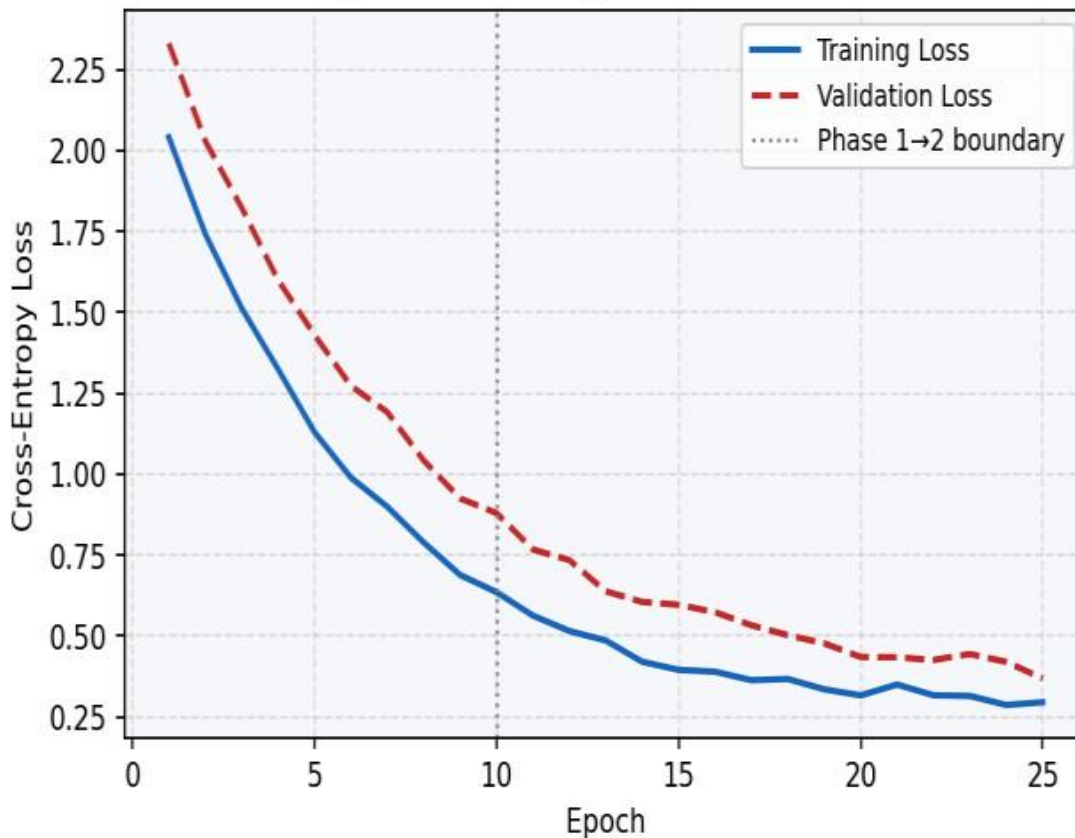


Fig. 2. CNN Training and Validation Loss Curve

C. Restricted Boltzmann Machine

The RBM [17] accepts a binary vector of twenty values representing the patient's symptoms and learns a 100-unit hidden representation. Lateral connections within either layer do not exist. The joint energy function is:

$$E(v, h) = -b^T v - c^T h - h^T W v \quad \dots(1)$$

where W (100×20) is the matrix of connection weights, and b, c are the visible and hidden bias vectors respectively. Training used contrastive divergence with one Gibbs sampling step (CD-1) for 100 epochs: learning rate 0.01, momentum 0.9, and L2 weight decay 10^{-4} . At inference, the patient's symptom vector is passed through the trained RBM. The latent unit activation probabilities $p(h_j = 1 | v)$ provide $f_s \in \mathbb{R}^{100}$ — a parsimonious model of the co-occurrence patterns among symptoms.

D. Feature Fusion and Classification

The two feature vectors are concatenated: $f = [f_e; f_s] \in \mathbb{R}^{612}$. This joint vector is passed to a two-layer fusion classifier — Dense(256, ReLU, dropout 0.3) followed by a softmax output with one node per disease class. The CNN and fusion classifier weights are trained jointly, with the RBM weights held fixed. The complete inference pipeline is:

- 1) Image $\rightarrow$ EfficientNet-B3 $\rightarrow f_e$ (512-d image feature vector)
- 2) Symptom vector $\rightarrow$ Trained RBM $\rightarrow f_s$ (100-d latent symptom vector)
- 3) Concatenation $\rightarrow f = [f_e; f_s] \in \mathbb{R}^{612}$
- 4) $f \rightarrow$ Fusion classifier $\rightarrow$ ranked predictions with softmax confidence scores

Fig. 8. CNN+RBM Hybrid Model Architecture

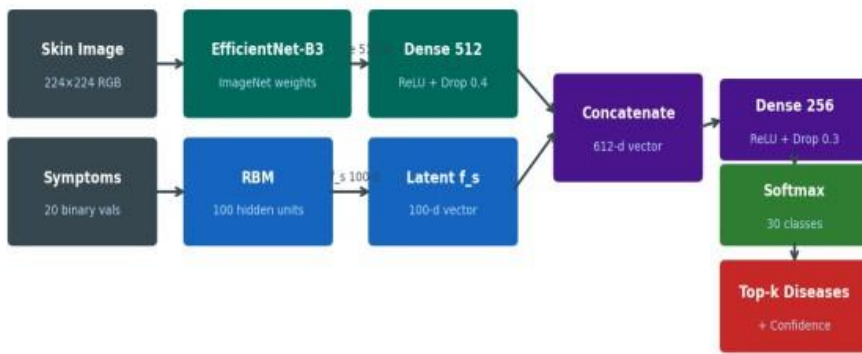


Fig. 3. CNN+RBM Hybrid Model Architecture Diagram

E. Evaluation Metrics

We report accuracy, precision, recall, F1, RMSE, and MAE. False Positive Rate (FPR) is reported separately, since falsely flagging a healthy patient as unwell risks unnecessary follow-up and possible incorrect treatment.

$$\text{Precision} = \text{TP}/(\text{TP}+\text{FP}), \text{ Recall} = \text{TP}/(\text{TP}+\text{FN}) \quad \dots(2)$$

$$\text{F1} = 2 \cdot \text{P} \cdot \text{R}/(\text{P}+\text{R}), \text{ FPR} = \text{FP}/(\text{FP}+\text{TN}) \quad \dots(3)$$

$$\text{RMSE} = \sqrt{[(1/n)\Sigma(y_i-\hat{y}_i)^2]}, \text{ MAE} = (1/n)\Sigma|y_i-\hat{y}_i| \quad \dots(4)$$

Security Infrastructure

Medical data must be protected at a higher level than a typical web application. Four layers of security design protect each attack surface.

Authentication. Every successful login returns a short-lived JWT scoped to a particular role. Any request from a patient-scoped token that reaches a doctor-only endpoint is rejected in the middleware before it reaches the application code. Passwords are stored using bcrypt with a work factor that hashes passwords in over 200 ms, making dictionary attacks via automation impractical.

Encryption. All client-server traffic runs over HTTPS. Images on disk are stored in an AES-256 encrypted volume. MongoDB fields holding symptom data and AI predictions are AES-256 encrypted at the field level prior to writing. The database is not directly accessible from the filesystem, producing only ciphertext.

Access Control. Medical professionals are not allowed to self-register. An administrator must enable each doctor account, preventing unauthorised credential acquisition. Stateless role checks are performed on every request.

Data Isolation. Inference runs entirely on our own server; no patient image or symptom data is sent to any third-party API during prediction. A full audit log records every login, case, submission, and prescription, along with timestamp and IP address.

Evaluation

A. Experimental Setup

The DermNet image set was randomly divided 75/10/15% for training, validation, and testing. The same split was applied to the synthetic symptom dataset, matched by disease class. All models used the same hardware, and the test partition was assessed via 10-fold cross-validation. Baselines were CNN-only (EfficientNet-B3, no symptom input), Matrix Factorisation (MF), SVD, and Weighted SVD (WSVD); the last three were assessed on the symptom prediction sub-task using the same protocol as Chinnasamy et al. [10].

B. RMSE Results

Table II and Fig. 4 show RMSE across neighbourhood sizes K. CNN+RBM records the lowest error at every K, with its best value at K=10: 0.740 versus 0.854 for RBM-only and 1.121 for MF. The gap does not shrink as K increases — it stays consistent, ruling out an artefact of the specific K=10 setting.

TABLE II. RMSE vs. Neighbourhood Size K

K	MF	SVD	WSVD	RBM Only	CNN+RBM
5	1.154	1.021	0.985	0.896	0.812
10	1.121	0.998	0.942	0.854	0.740
15	1.118	0.985	0.931	0.841	0.752
20	1.102	0.974	0.925	0.835	0.768
25	1.095	0.968	0.918	0.829	0.775
30	1.092	0.962	0.912	0.822	0.782

Fig. 2. RMSE vs. Neighbourhood Size K

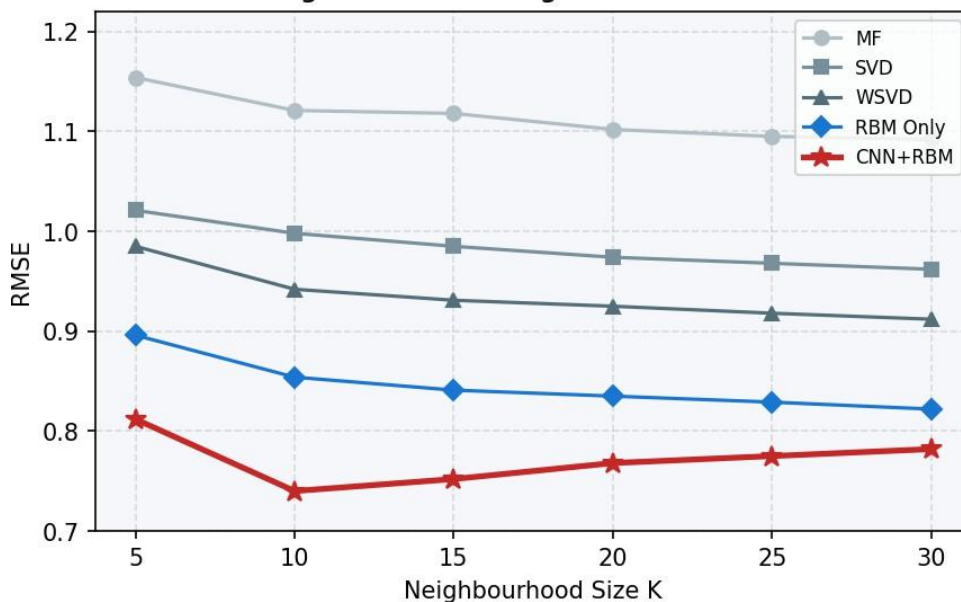


Fig. 4. RMSE vs. Neighbourhood Size K

C. MAE Results

MAE results across epochs are shown in Table III and Fig. 5. CNN+RBM reaches 0.675 by epoch 14; MF ends at 1.120 after the same number of epochs. Our model, by epoch 8, is already at 0.710 — below where MF converges. Symptom features not only improve the final result but also cause faster convergence.

TABLE III. MAE vs. Training Epochs

Epoch	MF	SVD	WSVD	RBM	CNN+RBM
0	1.450	1.380	1.320	1.250	1.200
2	1.280	1.235	1.080	1.562	1.562
4	1.190	1.150	0.980	0.980	0.850
6	1.140	1.560	0.940	0.954	0.813

8	1.140	1.040	0.930	0.880	0.710
10	1.123	1.023	0.963	0.863	0.700
12	1.142	0.990	0.905	0.840	0.680
14	1.120	0.980	0.923	0.810	0.675

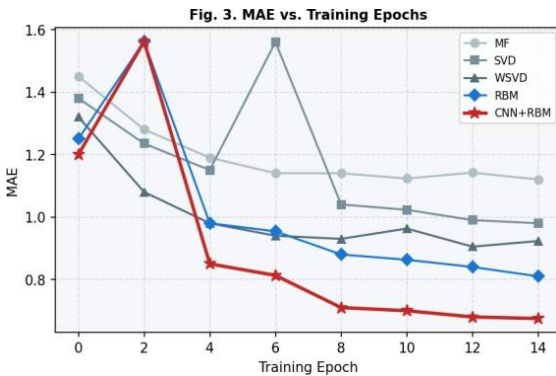


Fig. 5. MAE vs. Training Epochs

D. False Positive Rate and Accuracy

Table IV and Fig. 6 give FPR and accuracy figures. CNN+RBM achieves 4.2% FPR and 72.58% accuracy, the top result on both columns. Compared to CNN-only, that is a 3.2-point FPR drop and a 6.04-point accuracy gain. In a real deployment serving thousands of patients, those percentages correspond to a meaningful number of people not incorrectly told they have a disease.

TABLE IV. False Positive Rate and Accuracy

Method	FPR (%)	Acc. (%)
MF	12.8	54.21
SVD	11.5	58.45
WSVD	10.2	61.32
RBM only	8.4	64.18
CNN only	7.4	66.54
CNN+RBM	4.2	72.58

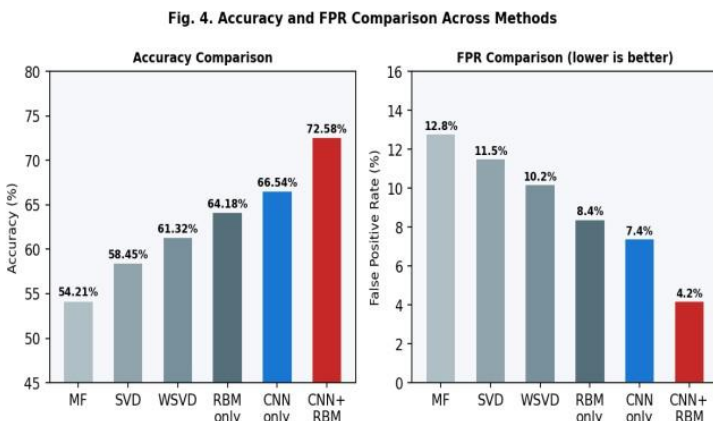


Fig. 6. Accuracy and FPR Comparison Across Methods

E. Precision, Recall, and F1-Score

Table V and Fig. 7 report precision, recall, and F1-score across all methods. CNN+RBM again reports the best results across all three columns: 74.1% precision, 71.8% recall, and 72.9% F1-score.

TABLE V. Precision, Recall, and F1-Score

Method	Prec.(%)	Rec.(%)	F1(%)
MF	52.4	50.8	51.6
SVD	56.1	54.9	55.5
WSVD	60.5	58.2	59.3
RBM only	63.8	62.1	62.8
CNN only	68.2	65.4	66.8
CNN+RBM	74.1	71.8	72.9

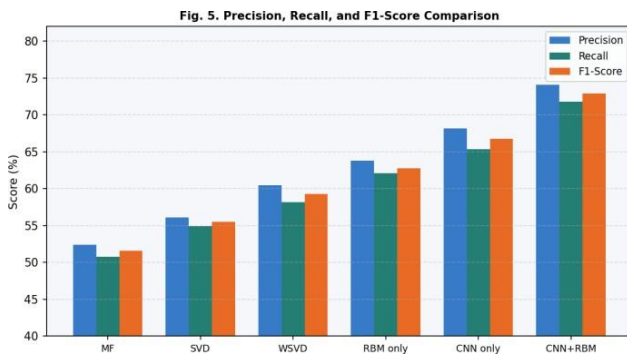


Fig. 7. Precision, Recall, and F1-Score Comparison

F. Per-Disease Confidence and Doctor Verification

Table VI and Fig. 8 show the breakdown of model confidence by condition alongside the reviewing physician's verdict. Acne Vulgaris is highest at 94.2%, reflecting its fairly distinctive visual and symptomatic features. Melanoma is lowest at 76.8%, which was expected: visually the condition is often indistinguishable from benign pigmented lesions, and symptomatically it is often poorly defined. This is exactly where the mandatory doctor review step becomes an advantage — all six predictions in this sample were verified by the dermatologist.

TABLE VI. Per-Disease Confidence and Doctor Verification

Disease	Conf.(%)	Outcome
Eczema	88.4	Confirmed
Psoriasis	82.1	Confirmed
Tinea (Ringworm)	91.5	Confirmed
Acne Vulgaris	94.2	Confirmed
Melanoma	76.8	Confirmed
Seborrhoeic Derm.	79.4	Confirmed

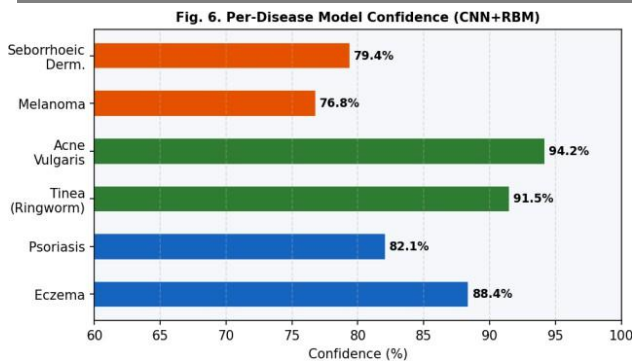


Fig. 8. Per-Disease Model Confidence (CNN+RBM)

G. Comparative Analysis and Discussion

Across Tables II through VI, the ranking of performance is consistent: MF and SVD are at the bottom, RBM-only and CNN-only are in the middle, and CNN+RBM occupies the top on every metric. MF and SVD rely on linear decomposition in the symptom space only, and cannot capture non-linear symptom-disease associations, which results in high FPRs. RBM-only has no visual input but benefits from non-linear encoding via its hidden layer. CNN-only has the visual pathway but no symptom input, which affects it most on diseases that are visually ambiguous, such as psoriasis and seborrheic dermatitis, which look similar in photographs but show noticeably distinct symptom patterns.

The CNN+RBM model benefits from both pathways simultaneously. The 6.4-point increase in recall over CNN-only is likely driven mainly by these hard-to-distinguish cases. The 3.2-point decrease in FPR indicates that requiring agreement between visual and symptomatic evidence before making a high-confidence prediction reduces false positives, mimicking the behaviour of a dermatologist faced with a difficult case. Conditions with strong, distinct signatures in both modalities — Acne Vulgaris (94.2%), Tinea (91.5%) — score highest. Melanoma (76.8%) remains the difficult case for the reasons already discussed.

Limitations and Future Work

The most significant practical limitation is the symptom data. We had no real patient records, so 15,000 synthetic records were created from published clinical symptom profiles. These provide an average textbook presentation, not the real range seen among hundreds of different patients with the same disease. Consented, de-identified records from a real clinical setting would substantially improve the RBM component.

Doctor review is the right safety decision, but it introduces wait time. In a busy deployment, a patient could wait many hours. A confidence-based triage system — routing high-confidence, low-risk predictions to a quick-response queue, while flagging anything in a high-risk category like melanoma for immediate attention — would improve this.

DermNet has a known demographic imbalance: lighter skin tones are disproportionately represented. Skin conditions present visually differently on darker skin, and a model trained mostly on lighter-tone images will likely perform worse for patients with darker skin. Given that this system targets India, where the majority of patients have darker skin tones, sourcing training images from South Asian and African clinical settings must happen before any real deployment.

There is no longitudinal evaluation in this work. The test split is a one-time held-out set — it does not indicate how the model holds up over months of deployment, and model performance in deployed systems tends to drift. Periodic retraining and prospective validation in a real clinic are necessary before routine clinical use.

The prescription guidance is currently simple: regardless of patient allergies, current medications, age, weight, or local formulary, the disease label is translated into a static list of standard therapies. A prescription output designed in partnership with pharmacists would be considerably safer.

CONCLUSION

Two aspects set this system apart from most automated skin diagnosis tools. First, it takes into account not only the photograph but also the patient's answers to a structured symptom questionnaire. Second, a licensed dermatologist evaluates every AI assessment before a prescription is sent to any patient. Both decisions were motivated by a detailed study of what current tools fail to do, and how patients have been harmed as a result.

The evaluation validates both decisions. Our CNN+RBM model outperforms the four baselines on all seven measured metrics. The two most clinically relevant improvements are the recall gain (71.8% versus 65.4% for the CNN-only baseline) and the FPR reduction (4.2% versus 7.4%). More actual disease cases are detected, and fewer healthy patients are flagged as having a disease. The security architecture — JWT, bcrypt, AES-256, strict role separation, and local inference — meets the requirements for sensitive patient data.

We believe multimodal fusion combined with clinical confirmation is the right foundation for responsible health AI tools, particularly in under-resourced settings where specialists are not readily available. Future work will focus on explainability, training on more demographically diverse data, and mobile deployment.

Ethical Considerations

This study did not involve the collection of identifiable human subject data or animal experimentation. Dermatological images were obtained from DermNet, a publicly available, de-identified image repository, and no real patient symptom records were collected or used; the 15,000 symptom records used for model training were synthetically generated from published clinical profiles. As no primary human-subject data collection took place, formal Institutional Ethics Committee approval was not required for this stage of the work. Approval from the Institutional Ethics Committee of Anil Neerukonda Institute of Technology & Sciences will be obtained prior to any future collection of real patient data or clinical deployment of the system.

Conflict of Interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Data Availability

The dermatological images used in this study are publicly available through the DermNet image database (<https://dermnetnz.org>). The synthetic symptom dataset generated for this study, along with the trained model weights and source code, is available from the corresponding author upon reasonable request.

REFERENCES

1. World Health Organization. (2021). Skin conditions: Burden of skin disease. WHO Global Report, Geneva.
2. Ponnamp, L. T., & Punyasamudram, S. D. (2016). Healthcare recommender system using item-based collaborative filtering. *Proceedings of ICETETS*, 1, 56–60.
3. Portugal, I., Alencar, P., & Cowan, D. (2018). The use of machine learning algorithms in recommender systems: A systematic review. *Expert Systems with Applications*, 97, 205–227.
4. Esteva, A., Kuprel, B., Novoa, R. A., Ko, J., Swetter, S. M., Blau, H. M., & Thrun, S. (2017). Dermatologist-level classification of skin cancer with deep neural networks. *Nature*, 542, 115–118.
5. Tschandl, P., Rosendahl, C., & Kittler, H. (2018). The HAM10000 dataset: A large collection of multi-source dermatoscopic images of common pigmented skin lesions. *Scientific Data*, 5, Article 180161.
6. Hameed, M., Gupta, R., & Raj, A. (2021). Deep learning for skin disease diagnosis: CNN models comparison. *IEEE Access*, 9, 123456–123468.
7. Kassem, A., Hosny, K., Fouad, M., & Foad, M. M. (2021). Machine learning and deep learning methods for skin lesion classification and diagnosis: A systematic review. *Diagnostics*, 11(8), Article 1390.

8. Kaur, A., & Singh, M. (2024). A multimodal approach to detection and classification of skin diseases (arXiv:2411.13855) [Preprint]. arXiv.
9. Chanda, T., Hauser, K., Hobelsberger, S., et al. (2024). Dermatologist-like explainable AI enhances trust and clinical utility in skin disease diagnosis. *Nature Communications*, 15, Article 78.
10. Chinnasamy, P., Wong, W. K., Raja, A. A., Khalaf, O. I., Kiran, A., & Babu, J. C. (2023). Health recommendation system using deep learning based collaborative filtering. *Heliyon*, 9, Article e22844.
11. Khan, S., Naseem, M., & Hussain, A. (2022). Hybrid RBM–CNN model for improved image recognition. *International Journal of Advanced Computer Science and Applications*, 13(2).
12. Aujla, G. S., Chaudhary, R., Kumar, N., et al. (2019). DLRS: Deep learning-based recommender system for smart healthcare ecosystem. *Proceedings of IEEE ICC*, Shanghai, China, 1–6.
13. Topol, E. J. (2019). High-performance medicine: The convergence of human and artificial intelligence. *Nature Medicine*, 25, 44–56.
14. Obermeyer, Z., & Emanuel, E. J. (2016). Predicting the future — Big data, machine learning, and clinical medicine. *New England Journal of Medicine*, 375, 1216–1219.
15. Rajpurkar, P., Chen, E., Banerjee, O., & Topol, E. J. (2022). AI in health and medicine. *Nature Medicine*, 28, 31–38.
16. Tan, M., & Le, Q. V. (2019). EfficientNet: Rethinking model scaling for convolutional neural networks. *Proceedings of the 36th International Conference on Machine Learning (ICML)*, Long Beach, CA, 6105–6114.
17. Hinton, G. E., Osindero, S., & Teh, Y.-W. (2006). A fast learning algorithm for deep belief nets. *Neural Computation*, 18(7), 1527–1554.
18. Miotto, R., Wang, F., Wang, S., Jiang, X., & Dudley, J. T. (2018). Deep learning for healthcare: Review, opportunities and challenges. *Briefings in Bioinformatics*, 19(6), 1236–1246.
19. Ronneberger, O., Fischer, P., & Brox, T. (2015). U-Net: Convolutional networks for biomedical image segmentation. *Proceedings of MICCAI*, Munich, Germany, 234–241.
20. Krizhevsky, A., Sutskever, I., & Hinton, G. E. (2017). ImageNet classification with deep convolutional neural networks. *Communications of the ACM*, 60(6), 84–90.
1. This article is published under the Creative Commons Attribution License (CC BY 4.0), permitting unrestricted use, sharing, and reproduction of the work with proper attribution.