

Original Research**Physician Perspectives on Generative AI in the Era of Healthcare Digitalisation:
Impact on Productivity, Resource Optimisation, and Cost Efficiency in India**Hitesh Kumar¹, Ritvik Prakash², Saurav Prakash³, Amiya Mrinal⁴¹Senior Resident, Dept of Surgical Oncology, Rajendra institute of medical sciences, Ranchi²Senior Resident, Dept of Plastic Surgery, Rajendra institute of medical sciences, Ranchi³PGP (2026), Indian Institute of Management Shillong, Meghalaya, India⁴PGP (2025), Indian Institute of Management Shillong, Meghalaya, India

ARTICLE INFO



Keywords: Generative AI; ChatGPT;
Physicians; Healthcare Digitalisation;
Documentation Burden; Productivity;
India

ABSTRACT

Background: Generative artificial intelligence (GenAI) tools such as ChatGPT, Gemini, and Copilot are rapidly entering clinical workflows worldwide. Their real-world reception among Indian physicians — who work under distinctive resource, cost, and volume pressures — remains under-explored.

Objective: To assess Indian physicians' perceptions of GenAI in terms of clinical productivity, documentation burden, administrative workload, information access, cost efficiency, and readiness for routine integration.

Methods: A cross-sectional online survey was conducted among 53 practising physicians across government, private, corporate, and academic settings in India. The instrument comprised five demographic items, one usage-frequency item, nine five-point Likert statements, and one open-ended question. Descriptive statistics, Cronbach's alpha, Spearman correlation, Kruskal–Wallis, Mann–Whitney U, and chi-square tests were applied at $\alpha = 0.05$.

Results: Daily GenAI users accounted for 45.3% of respondents; only 3.8% were non-users. Physicians reported strongest agreement that GenAI reduces documentation time (mean 4.19 ± 0.71) and improves medical writing quality (4.19 ± 0.79), followed by administrative workload reduction (4.09 ± 0.63) and clinical information access (4.04 ± 0.65). Perceived impact on patient-facing time (3.38 ± 0.93) and operational cost reduction (3.59 ± 0.80) was more modest. Internal consistency was good (Cronbach's $\alpha = 0.80$). Frequency of GenAI use correlated positively with future-use intention ($p = 0.415$, $p = 0.002$), integration readiness ($p = 0.328$, $p = 0.016$), medical writing benefit ($p = 0.359$, $p = 0.008$), productivity ($p = 0.302$, $p = 0.028$), and documentation efficiency ($p = 0.282$, $p = 0.041$). Attitudes did not differ significantly by qualification, practice setting, age, or experience (all $p > 0.05$).

Conclusion: Indian physicians hold predominantly favourable views of GenAI, particularly for documentation and knowledge-work efficiency, with more cautious appraisals of its economic and patient-time dividends. Higher exposure is linked to greater acceptance, supporting structured training, ethical governance, and pilot integration into routine clinical practice.

Introduction

Healthcare in India is undergoing a compressed digital transition. The Ayushman Bharat Digital Mission, expanding electronic health record (EHR) adoption, and rising outpatient volumes have made the interface between clinicians and digital tools a decisive determinant of care quality [20,22]. Simultaneously, the public release of large language models (LLMs) — ChatGPT in November 2022, followed by Gemini, Copilot, and Claude — has moved artificial intelligence (AI) from radiology labs into everyday physician workflows for note-writing, patient education, literature review, and administrative correspondence [1,2,14,16].

Documentation and administrative load are consistently identified as

the largest non-clinical drains on physician time and among the strongest predictors of burnout. Ambient AI and generative documentation tools have already been shown to shorten note times, decrease after-hours charting, and reduce self-reported burnout in ambulatory clinicians in the United States [3,4,5]. Health-system leaders in the US now rank generative AI for clinical documentation as the single most widely deployed AI use case, ahead of even diagnostic imaging [6]. A 2024 American Medical Association survey reported that 66% of US physicians used health-AI in 2024, a 78% increase over 2023, driven largely by generative tools [7].

Physician attitudes, however, are shaped by more than technical

* Corresponding author: Dr. Saurav Prakash, Indian Institute of Management, Shillong. Email: sauravprakash121@gmail.com

capability. Recent international cross-sectional studies show that acceptance is predicted by AI literacy, ethical sensitivity, perceived usefulness, and organisational readiness [8,9,12]. Concerns about hallucination, accountability, data privacy, medico-legal exposure, and the risk of clinician deskilling remain material and unresolved [2,10,17,21,24].

Indian evidence, however, is thin. Most published work is narrative or opinion-based; primary survey data from Indian physicians on GenAI's effects on productivity, resource optimisation, and cost efficiency are scarce. Given India's constrained per-capita clinician availability, cost sensitivity, and documentation-heavy public sector, physician perceptions here may diverge meaningfully from Western datasets. This study addresses that gap.

Objectives: (i) to quantify Indian physicians' agreement with GenAI's impact on productivity, documentation, patient time, medical writing, administrative workload, information access, cost, and integration; (ii) to test whether frequency of GenAI use predicts more favourable perceptions; (iii) to examine whether perceptions vary by qualification, practice setting, age, or experience; and (iv) to extract qualitative themes on the greatest benefit and biggest challenge.

Materials and Methods

Study Design and Setting

A cross-sectional, questionnaire-based, observational study was conducted online among Indian physicians during the survey period recorded in the response log. Reporting follows the Strengthening the Reporting of Observational Studies in Epidemiology (STROBE) guideline (Appendix A). Figure 1 presents the participant flow diagram.

Participants

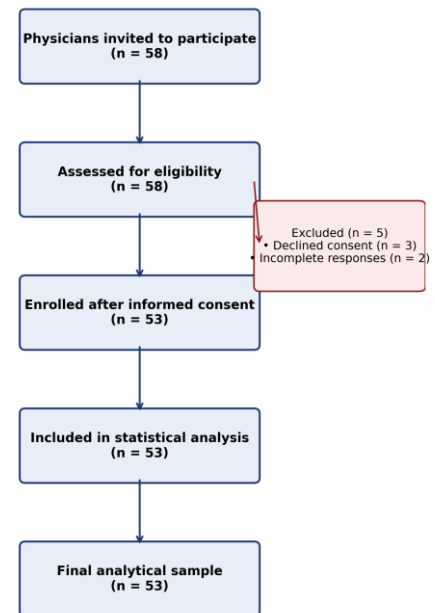
Inclusion criteria: MBBS, BDS, MD/MS/DNB, MDS, and DM/MCh qualified physicians currently in clinical practice in India who provided informed consent. Exclusion criteria: non-clinical roles, students without a primary medical qualification, and incomplete consent. Convenience and snowball sampling were used via professional networks. The final analysable sample comprised 53 respondents.

Instrument

The self-designed instrument, administered through Google Forms, included: an informed-consent statement; demographic items (age group, medical qualification, years of clinical experience, type of practice); frequency of GenAI use (Never/Rarely/Monthly/Weekly/Daily); nine five-point Likert items

(Strongly Disagree – Strongly Agree) covering overall clinical productivity, documentation time, patient time, medical writing quality, administrative workload, clinical information access, operational cost, integration readiness, and future-use intention; and one open-ended question on the perceived greatest benefit or biggest challenge. The complete questionnaire is provided in Supplementary Material

Figure 1. STROBE Participant Flow Diagram



Ethical Considerations

Participation was voluntary and informed consent was mandatory. Responses were anonymised and no personal identifiers were used in analysis. The study followed the principles of the Declaration of Helsinki.

Statistical Analysis

Analyses were performed in Python (pandas, SciPy) with outputs verified against standard SPSS-equivalent computations. Categorical data are presented as frequencies and percentages; Likert responses were coded 1–5 and summarised using mean, standard deviation (SD), median, and the proportion responding Agree or Strongly Agree. Internal consistency of the nine-item attitude scale was assessed with Cronbach's alpha. Spearman's rank-order correlation examined associations between frequency of GenAI use and each item. Kruskal–Wallis H tests compared composite attitude scores across qualification, practice setting, age, and experience subgroups. Mann–Whitney U compared Under-30 versus 30–39 year respondents. Chi-square evaluated the association between daily use and future-use intention. Statistical significance was set at $p < 0.05$ (two-tailed). Open-ended responses were analysed by inductive thematic coding.

Results

Sample Characteristics

Fifty-three physicians completed the survey. Table 1 shows the age and experience distribution; Table 2 presents medical qualification and practice setting. The sample was young (71.7% under 30 years) and predominantly early-career (77.4% with less than five years of practice). Table 3 summarises the frequency of GenAI use.

Table 1. Demographic Characteristics of Respondents (N = 53).

Variable	Category	n	%
Age Group	Under 30 years	38	71.7
	30 to 39 years	15	28.3
Years of Clinical Experience	Less than 5 years	41	77.4
	5 to 10 years	10	18.9
	11 to 20 years	2	3.8

n, number of respondents. Percentages are rounded to one decimal place.

Table 2. Participant Characteristics: Medical Qualification and Practice Setting (N = 53).

Variable	Category	n	%
Medical Qualification	BDS	23	43.4
	MD / MS / DNB	13	24.5
	MBBS	7	13.2
	DM / MCh	6	11.3
	MDS	4	7.5
Practice Setting	Government Hospital	27	50.9
	Medical College	11	20.8
	Private Hospital	7	13.2
	Corporate Hospital	4	7.5
	Private Clinic	3	5.7
	Other	1	1.9

BDS, Bachelor of Dental Surgery; MBBS, Bachelor of Medicine, Bachelor of Surgery; MDS, Master of Dental Surgery; MD, Doctor of Medicine; MS, Master of Surgery; DNB, Diplomate of National Board; DM, Doctor of Medicine (super-specialisation); MCh, Master of Chirurgiae.

Table 3. Frequency of Generative AI Use for Professional Purposes (N = 53).

Frequency	n	%
Daily	24	45.3
Weekly	16	30.2
Monthly	6	11.3
Rarely	5	9.4
Never	2	3.8

Daily or weekly use combined accounted for 75.5% of respondents.

Table 4. Productivity Outcomes: Descriptive Statistics for Perceived Impact of Generative AI on Clinical Productivity (N = 53).

Item	Statement	Mean	SD	Median	% Agree/SA
Q1	Improves overall clinical productivity	3.98	0.67	4	86.8
Q2	Reduces time spent on documentation	4.19	0.71	4	86.8
Q3	Allows more time with patients	3.38	0.93	4	50.9

Scale: 1 = Strongly Disagree; 5 = Strongly Agree. SD, standard deviation; SA, Strongly Agree.

Table 5. Information Quality and Communication: Perceived Impact of Generative AI (N = 53).

Item	Statement	Mean	SD	Median	% Agree/SA
Q4	Improves quality of medical writing / educational material	4.19	0.79	4	90.6
Q5	Reduces administrative workload	4.09	0.63	4	88.7
Q6	Improves access to relevant clinical information	4.04	0.65	4	84.9

Scale: 1 = Strongly Disagree; 5 = Strongly Agree.

Table 6. Cost Efficiency: Perceived Impact of Generative AI on Operational Costs (N = 53).

Item	Statement	Mean	SD	Median	% Agree/SA
Q7	Can reduce operational costs in healthcare	3.59	0.80	4	56.6

Cost reduction had the second-lowest endorsement in the scale.

Table 7. Thematic Analysis of Open-Ended Responses (N = 53).

Theme	Frequency of Mention	Illustrative Quotation
Efficiency / Time Saving	14	"AI has reduced clerical workload tremendously."
Accuracy / Hallucination	9	"Generative AI is never accountable for its work and produces false statements frequently."
Information Access	8	"Real-time help with quick information leading to improved decision making."
Medical Writing Support	6	"It reduces my time on sentence framing, improves my writing work."
Deskilling / Overreliance	6	"Overdependence on AI reduces clinicians' thinking ability."
Ethics / Privacy / Data Security	5	"Ethical issues and plagiarism, data security is questionable."
Empathy / Workforce Displacement	3	"Practising medicine requires a lot of empathy."
Patient-facing Risks	3	"Patients using AI at home as medical counsel."

Frequencies reflect the number of respondents who explicitly referenced the theme. Quotations are verbatim; minor grammatical edits for readability.

Table 8. Future Adoption: Integration Readiness and Intention (N = 53).

Item	Statement	Mean	SD	Median	% Agree/SA
Q8	Should be integrated into routine clinical practice	3.89	0.67	4	75.5
Q9	Intend to increase future use	4.09	0.56	4	88.7

Future-use intention was the item most strongly associated with current frequency of use ($\rho = 0.415$, $p = 0.002$).

Table 9 presents reliability statistics including item–total correlations and Cronbach's α if the item is deleted. All items contributed positively to scale reliability.

Table 9. Reliability Analysis of the 9-Item Attitude Scale (N = 53).

Item	Corrected Item–Total Correlation	Cronbach's α if Item Deleted
Q1	0.421	0.789
Q2	0.502	0.779
Q3	0.539	0.775
Q4	0.477	0.782
Q5	0.533	0.776
Q6	0.489	0.781
Q7	0.473	0.783
Q8	0.545	0.774
Q9	0.482	0.783
Overall Cronbach's α	—	0.800

Cronbach's $\alpha \geq 0.70$ is considered acceptable; ≥ 0.80 good; ≥ 0.90 excellent.

Table 10. Correlation and Subgroup Analysis: Frequency of Use, Composite Attitude, and Demographic Comparisons (N = 53).

Test	Statistic (ρ / H / U / χ^2)	P-value	Significance
Spearman: Q1 Productivity vs Use Frequency	0.302	0.028	*
Spearman: Q2 Documentation vs Use Frequency	0.282	0.041	*
Spearman: Q3 Patient time vs Use Frequency	0.033	0.813	ns
Spearman: Q4 Medical writing vs Use Frequency	0.359	0.008	**
Spearman: Q5 Admin workload vs Use Frequency	0.067	0.635	ns
Spearman: Q6 Info access vs Use Frequency	0.223	0.108	ns
Spearman: Q7 Cost reduction vs Use Frequency	0.010	0.942	ns
Spearman: Q8 Routine integration vs Use Frequency	0.328	0.016	*
Spearman: Q9 Future intent vs Use Frequency	0.415	0.002	**
Spearman: Composite Score vs Use Frequency	0.305	0.027	*
Kruskal–Wallis (Composite by Qualification, H)	0.42	0.980	ns
Kruskal–Wallis (Composite by Practice Setting, H)	2.06	0.725	ns
Kruskal–Wallis (Composite by Age Group, H)	0.38	0.539	ns

Kruskal–Wallis (Composite by Experience, H)	0.33	0.850	ns
Mann–Whitney U (Under 30 vs 30–39)	254.0	0.545	ns
Chi-square (Daily use vs Future intent Agree/SA)	1.12	0.289	ns

* $p < 0.05$; ** $p < 0.01$; ns, not significant. ρ , Spearman rank correlation coefficient; H, Kruskal–Wallis H statistic; U, Mann–Whitney U statistic; χ^2 , chi-square statistic.

Attitudes Toward Generative AI

Internal consistency of the nine-item attitude scale was good (Cronbach's $\alpha = 0.80$; Figure 10 and Table 9). Tables 4–6 and Table 8 present item-level descriptives grouped by construct; Figure 4 shows the distribution of Likert responses, and Figure 5 the mean scores with error bars.

Frequency of Use Versus Attitudes

Table 10 summarises Spearman correlations between frequency of GenAI use and each attitude item, together with subgroup comparisons. Figure 6 provides a correlation heatmap for the full item set. Physicians who use GenAI more frequently expressed significantly more positive perceptions of productivity, documentation, medical writing, integration readiness, and — most strongly — future-use intention. Perceptions of patient-time, administrative workload, information access, and cost were not associated with frequency, suggesting these are structurally rather than experientially determined.

Subgroup Comparisons

Kruskal–Wallis tests on the composite score showed no statistically significant differences across qualification ($H = 0.42$, $p = 0.98$), practice setting ($H = 2.06$, $p = 0.73$), age group ($H = 0.38$, $p = 0.54$), or experience ($H = 0.33$, $p = 0.85$). The Mann–Whitney U test between Under-30 ($M = 3.92$) and 30–39 ($M = 3.99$) year respondents was non-significant ($U = 254$, $p = 0.55$). The chi-square test between daily use and future-use intention (Agree/SA vs other) was non-significant ($\chi^2 = 1.12$, $p = 0.29$), although the direction favoured daily users (95.8% vs 82.8%).

Figure 2. Demographic Profile of Respondents (N = 53)

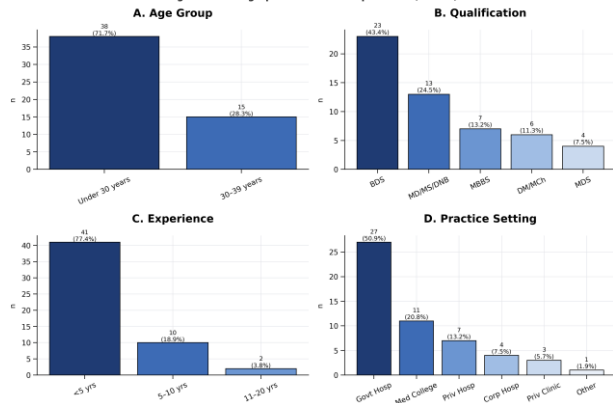


Figure 3. Frequency of Generative AI Use Among Respondents

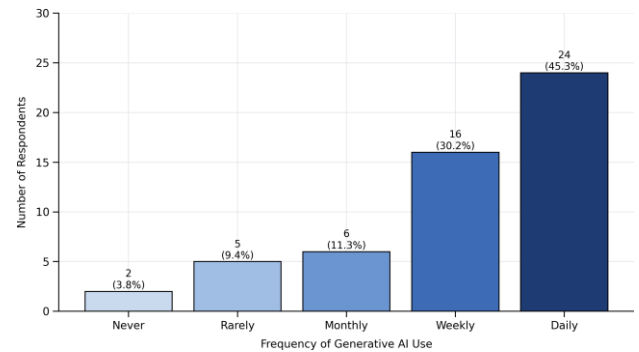


Figure 4. Distribution of Likert Responses Across Nine Attitude Items

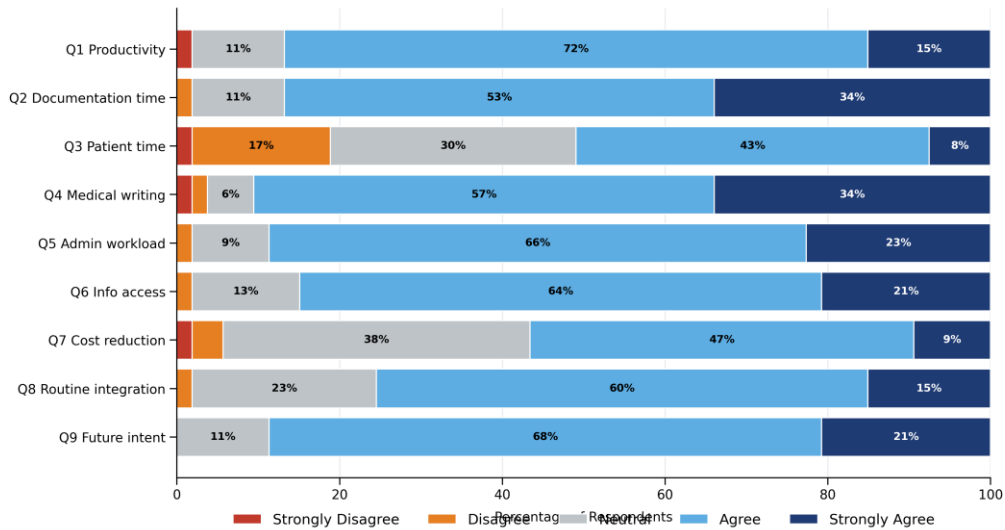


Figure 5. Mean Attitude Scores with Standard Deviations (N=53)

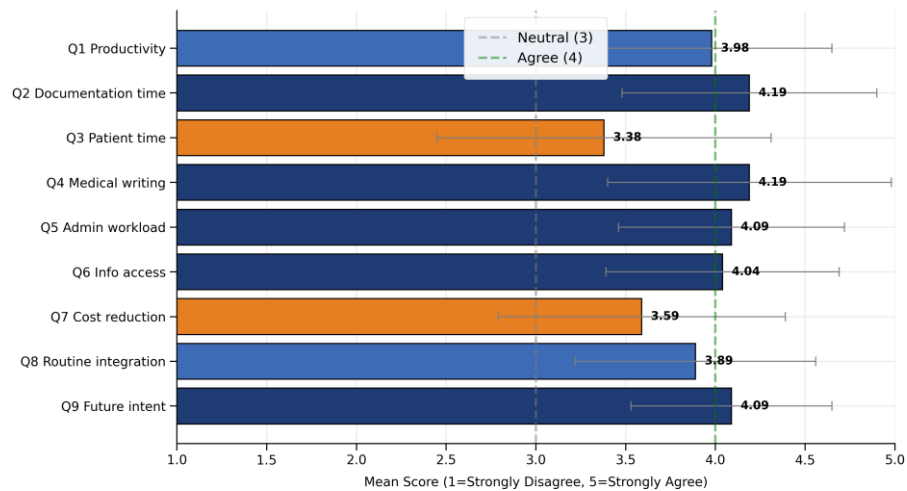


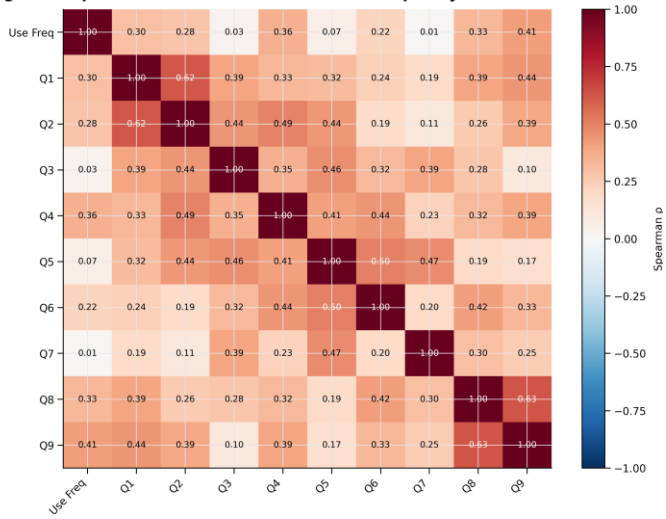
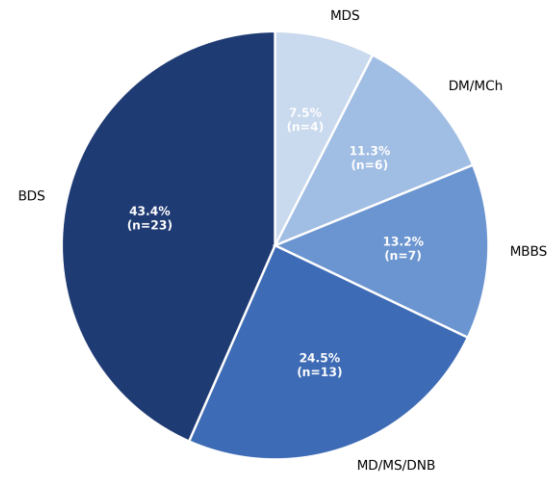
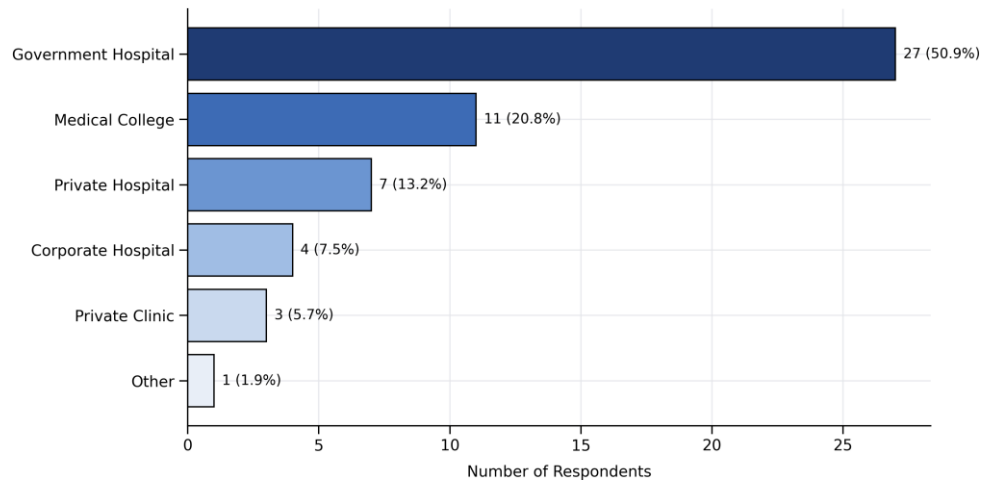
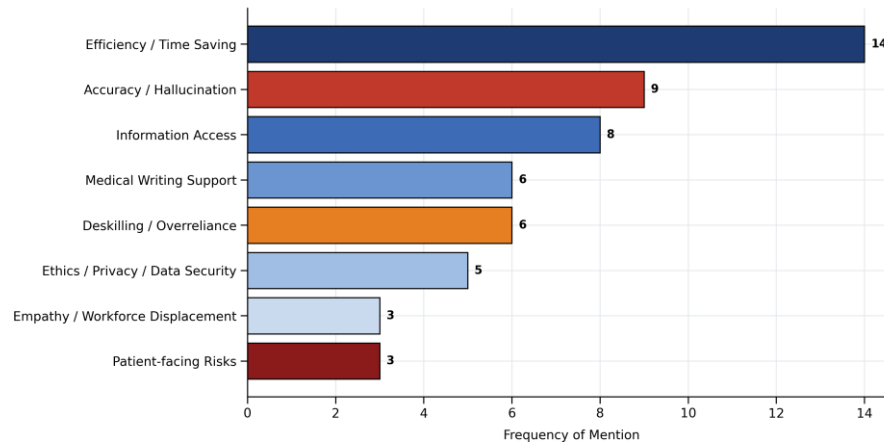
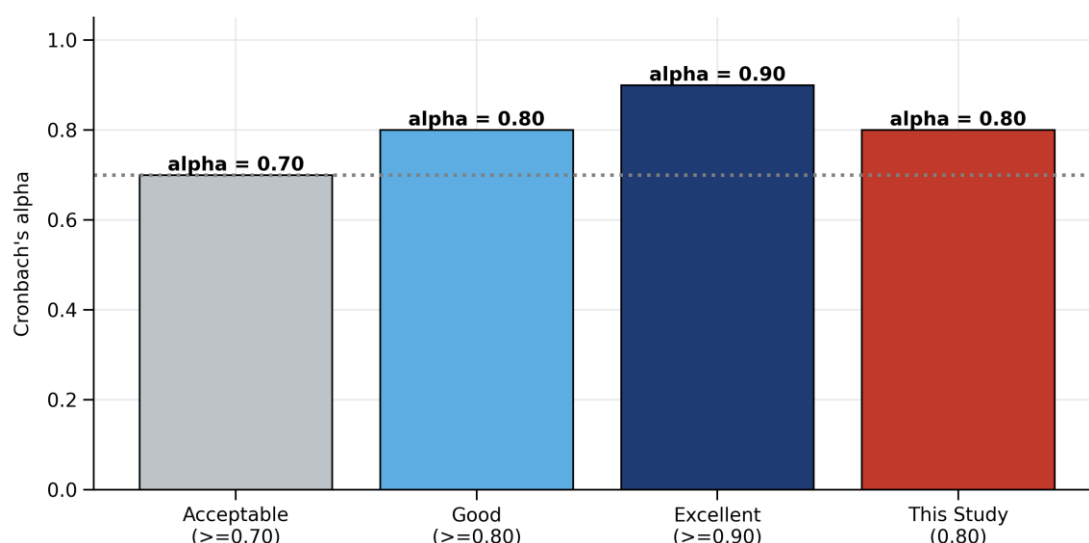
Figure 6. Spearman Correlation Matrix: Use Frequency and Attitude Items**Figure 7. Distribution of Medical Qualifications (N=53)****Figure 8. Distribution of Practice Settings (N=53)****Figure 9. Themes from Open-Ended Responses (Qualitative Analysis)**

Figure 10. Reliability of the 9-Item Attitude Scale (Cronbach's alpha = 0.80)



Thematic Analysis of Open-Ended Responses

Inductive coding of the open-ended item generated eight themes (Table 7 and Figure 9). Efficiency and time saving emerged as the dominant perceived benefit, followed by improved information access and support for medical writing. The most frequently expressed concern was accuracy and hallucination, followed by risks of clinician deskilling and unresolved questions on ethics and data privacy.

Discussion

Principal Findings

This study is among the first primary-data surveys quantifying Indian physicians' perceptions of GenAI across productivity, resource, and cost dimensions. Three findings stand out. First, uptake is already high: 75.5% of respondents use GenAI at least weekly, mirroring — and possibly outpacing — recent US rates where 66% of physicians reported any AI use in 2024 [7]. Second, perceived benefits cluster around cognitive-administrative work rather than economic or patient-facing outcomes. Documentation, medical writing, administrative workload, and information access all scored above 4.0 on the Likert scale, consistent with international evidence that ambient and generative AI's most reliable near-term gains are in note-writing and clerical offloading [3,4,5,11]. Third, exposure predicts acceptance: the composite score correlated significantly with frequency of use ($p = 0.305$, $p = 0.027$), and the strongest association was with future-use intention ($p = 0.415$, $p = 0.002$), echoing international findings that hands-on experience is a leading predictor of acceptance [8].

Contextualising Concerns

Qualitative themes on hallucination, accountability, and data privacy replicate concerns raised in AMA guidance for physicians [2], the ICMJE 2024 framework on AI in publications [10], and a 2025 physician survey where roughly half of respondents worried about GenAI increasing patient anxiety [9]. The Indian sample adds a distinctive economic worry — subscription creep and workforce displacement — that deserves policy attention as public-sector procurement scales up. This is consistent with EY and PwC analyses projecting GenAI-driven productivity gains of 30–32% in Indian healthcare by 2030 while cautioning about legacy-system integration costs.

Non-significant Subgroup Effects

The absence of significant differences by qualification, setting, age, or experience is itself informative. It suggests that GenAI perception is currently shaped less by clinical identity and more by tool familiarity — a pattern earlier reported for EHR acceptance [19]. It should be interpreted cautiously given modest cell sizes (e.g., only two respondents with 11–20 years of practice).

Limitations

The convenience sample of 53 physicians limits generalisability, especially to senior specialists (only 3.8% had more than 10 years of practice). Self-report bias and social-desirability effects are possible. The cross-sectional design cannot establish causality between use frequency and acceptance. The BDS-heavy sample (43.4%) partially

reflects the author's professional network. No objective outcome measure (e.g., documentation minutes, RVUs) was collected — a limitation shared with most international physician-attitude surveys and one that studies such as Husa et al. and Olson et al. have begun to close with EHR-metadata designs [3,4].

Implications for Policy and Practice

Four implications follow. First, structured training in prompt engineering, hallucination management, and PHI-safe use should be embedded in undergraduate and postgraduate curricula. Second, institutional GenAI-use policies are needed to standardise consent, data handling, and audit trails. Third, pilot ambient-documentation deployments in Indian tertiary settings — with pre/post EHR-metadata evaluation — would generate the first objective productivity evidence in this system. Fourth, formal economic evaluations are required to test whether GenAI is a net cost-saver or cost-shifter in the Indian public sector.

Conclusion

Indian physicians view Generative AI favourably, particularly as a documentation and knowledge-work accelerator, and intend to increase its use. Enthusiasm is tempered by well-founded concerns about accuracy, accountability, data privacy, cost creep, and clinician deskilling. Higher exposure predicts stronger acceptance, supporting a structured, ethically governed pathway for integrating GenAI into routine Indian clinical practice.

Declarations

Funding: None declared.

Conflict of Interest: The author declares no conflict of interest.

Data Availability: De-identified dataset available from the corresponding author upon reasonable request.

Ethics Approval: Voluntary participation with informed consent; anonymised data. The study followed the principles of the Declaration of Helsinki.

References

1. Bhuyan SS, Sateesh V, Mukul N, et al. Generative Artificial Intelligence Use in Healthcare: Opportunities for Clinical Excellence and Administrative Efficiency. *J Med Syst.* 2024;48. doi:10.1007/s10916-024-02136-1.
2. American Medical Association. ChatGPT and Generative AI: What Physicians Should Consider. Chicago: AMA; 2023.
3. Olson KD, Meeker D, Troup M, et al. Use of Ambient AI Scribes to

Reduce Administrative Burden and Professional Burnout. *JAMA Netw Open.* 2025;8(10):e2534976. doi:10.1001/jamanetworkopen.2025.34976.

4. Husa RA, Haggerty J, Nute AW, et al. Ambient Artificial Intelligence Use and Clinician Documentation Burden, Productivity, and Efficiency. *JAMA Netw Open.* 2026;9(5):e2615762. doi:10.1001/jamanetworkopen.2026.15762.

5. Albrecht M, Shanks D, Shah T, et al. Enhancing Clinical Documentation with Ambient Artificial Intelligence: a Quality-Improvement Survey. *JAMIA Open.* 2025;8(1):oaf013. doi:10.1093/jamiaopen/oaf013.

6. Poon EG, Lemak CH, Rojas JC, Guptill J, Classen D. Adoption of Artificial Intelligence in Healthcare: Survey of Health System Priorities, Successes, and Challenges. *J Am Med Inform Assoc.* 2025. doi:10.1093/jamia/ocaf065.

7. American Medical Association. Physician Sentiments Around the Use of AI in Health Care: 2024 Survey. Chicago: AMA; 2025.

8. Tosun N, Tosun M, Güner A. Physicians' Attitudes Toward Artificial Intelligence and ChatGPT: a Cross-Sectional Study. *J Health Organ Manag.* 2026. doi:10.1108/JHOM-12-2025-0889.

9. Faria V, Goturi N, Dynak AD, et al. Triadic Relations in Healthcare: Surveying Physicians' Perspectives on Generative AI Integration and Its Role on Empathy, the Placebo Effect and Patient Care. *Front Psychol.* 2025;16:1612215. doi:10.3389/fpsyg.2025.1612215.

10. Fakharifar A, Beizavi Z, Pouramini A, Haseli S. Application of Artificial Intelligence and ChatGPT in Medical Writing: a Narrative Review. *J Med Artif Intell.* 2025;8.

11. Topol EJ. High-Performance Medicine: The Convergence of Human and Artificial Intelligence. *Nat Med.* 2019;25(1):44–56.

12. Sallam M. ChatGPT Utility in Healthcare Education, Research, and Practice: Systematic Review on the Promising Perspectives and Valid Concerns. *Healthcare (Basel).* 2023;11(6):887.

13. Ayers JW, Poliak A, Dredze M, et al. Comparing Physician and AI Chatbot Responses to Patient Questions Posted to a Public Social Media Forum. *JAMA Intern Med.* 2023;183(6):589–596.

14. Lee P, Bubeck S, Petro J. Benefits, Limits, and Risks of GPT-4 as an AI Chatbot for Medicine. *N Engl J Med.* 2023;388:1233–1239.

15. Rajpurkar P, Chen E, Banerjee O, Topol EJ. AI in Health and Medicine. *Nat Med.* 2022;28:31–38.

16. Singhal K, Azizi S, Tu T, et al. Large Language Models Encode Clinical Knowledge. *Nature.* 2023;620:172–180.

17. Meskó B, Topol EJ. The Imperative for Regulatory Oversight of Large Language Models in Healthcare. *NPJ Digit Med.* 2023;6:120.

18. Shah NH, Entwistle D, Pfeffer MA. Creation and Adoption of Large Language Models in Medicine. *JAMA.* 2023;330:866–869.

19. Beam AL, Kohane IS. Big Data and Machine Learning in Health Care. *JAMA.* 2018;319:1317–1318.

20. Ministry of Health & Family Welfare, Government of India.

Ayushman Bharat Digital Mission — Health Data Management Policy. 2020.

21. Kelly CJ, Karthikesalingam A, Suleyman M, Corrado G, King D. Key Challenges for Delivering Clinical Impact with AI. *BMC Med.* 2019;17:195.
22. Ghosh A, Sengupta A. Artificial Intelligence in Indian Healthcare: Policy and Practice Review. *Indian J Public Health.* 2022;66(Suppl 1):S27–S33.
23. NITI Aayog. National Strategy for Artificial Intelligence. Government of India; 2018.
24. Panch T, Mattie H, Celi LA. The "Inconvenient Truth" About AI in Healthcare. *NPJ Digit Med.* 2019;2:77.
25. Esteva A, Robicquet A, Ramsundar B, et al. A Guide to Deep Learning in Healthcare. *Nat Med.* 2019;25:24–29.