

Voice-Controlled IoT Smart Notice Board for Educational Campuses

Dr. S. Kalaivani

Assistant Professor, Department of Bachelor of Computer Science, St Vincent Pallotti College,
Bengaluru, Karnataka, India

DOI: <https://doi.org/10.51244/IJRSI.2026.1307000329>

Received: 02 August 2026; Accepted: 07 August 2026; Published: 18 August 2026

ABSTRACT

Educational institutions continue to rely on conventional notice boards to disseminate information to students and staff, a process that involves manual preparation and posting of printed notices, resulting in delays, paper consumption, and limited accessibility for students who are off campus. This paper presents the design and a preliminary experimental evaluation of a low-cost, voice-controlled IoT notice board prototype for educational institutions. The proposed system uses an ESP32 microcontroller with a microphone module and the Google Speech-to-Text cloud API to convert spoken announcements into text, which is displayed locally on an LCD, written to a Firebase Cloud Database, and made available to students through a companion mobile application. A publisher review-and-confirm step is incorporated so that a transcribed notice must be approved before it is published, reducing the risk of transcription errors reaching students. The prototype was deployed in a single classroom setting for functional observation, and speech-recognition performance was measured on a defined 50-sentence test corpus read by five speakers under three ambient-noise conditions (quiet, moderate, and high noise, operationally defined by approximate decibel ranges). Mean word accuracy was 92% in quiet conditions, 87% under moderate noise, and 78% under high noise, with an unweighted three-condition average of approximately 85.7%. Under the trial protocol used, notice-publication time was reduced from a manual baseline of 8 to 10 minutes to approximately 45 seconds, a reduction of roughly 90 to 92 percent depending on the baseline used. The estimated one-time prototype hardware cost is approximately INR 2,000 to 2,500, exclusive of recurring cloud, connectivity, and maintenance costs. The results are preliminary and based on small-scale, single-site testing; they indicate promising directions for low-cost voice-driven campus communication rather than a fully validated institutional solution. Security, authentication, and privacy considerations that must be addressed before deployment are discussed explicitly, along with a research agenda for rigorous, larger-scale evaluation.

Keywords: Internet of Things, Smart Notice Board, Voice Recognition, ESP32, Firebase, Prototype Evaluation, Word Error Rate

INTRODUCTION

Rapid technological advancements have transformed educational institutions into digitally connected learning environments. Smart classrooms, online learning platforms, cloud-based attendance systems, and virtual laboratories have become common in higher education. However, one of the most widely used communication mechanisms — the traditional notice board — remains largely unchanged.

Conventional notice boards require administrative personnel or faculty members to prepare printed notices, physically attach them to display boards, and replace outdated notices periodically. This process is time-consuming and paper-intensive, and students who are absent from campus or unable to visit notice boards frequently may miss important announcements.

This paper reports the design of a voice-controlled IoT notice board prototype and a preliminary, small-scale experimental evaluation of its speech-recognition performance and notice-publication speed. Consistent with its exploratory scope, the study is framed as a prototype design and preliminary evaluation rather than a validated institutional solution; the methodology, results, and limitations sections define the evaluation conditions explicitly so that the reported figures can be interpreted and reproduced.

Problem Statement

Despite digital transformation in education, communication through notice boards continues to depend on manual procedures. Existing systems present several challenges: delay in publishing notices, high paper consumption, manual maintenance, students missing important announcements, lack of remote accessibility, and absence of a centralized digital archive. These limitations reduce communication efficiency and increase operational costs. Therefore, an automated, voice-enabled, IoT-based communication system is required — provided that authentication, review, and data-protection safeguards are designed in from the outset rather than treated as optional enhancements.

Research Objectives

The objectives of this study are to:

- Design a low-cost IoT-based notice board with a voice-to-text publishing workflow.
- Incorporate a publisher review-and-confirm step to reduce the risk of transcription errors being published unchecked.
- Measure speech-recognition performance under operationally defined noise conditions using Word Error Rate.
- Measure notice-publication time under a defined, repeatable protocol and report the calculation explicitly.
- Identify the security, privacy, and reliability requirements that must be satisfied before institutional deployment.
- Report the prototype's one-time hardware cost separately from anticipated recurring operating costs.

LITERATURE REVIEW

Internet of Things (IoT) systems connect everyday devices to the internet to exchange data with minimal human intervention, and have been applied broadly to smart-city and smart-campus infrastructure (Gubbi et al., 2013; Zanella et al., 2014; Al-Fuqaha et al., 2015). Within this space, digital notice boards have attracted sustained interest as a low-cost demonstration of IoT communication.

Early systems used GSM modules to transmit SMS-based updates directly to a microcontroller-driven display, avoiding manual rewriting but incurring recurring per-message costs. Bluetooth-based boards improved convenience for local updates but were constrained by short communication range. More recent work has shifted to Wi-Fi-enabled microcontrollers such as the ESP8266 and ESP32, which allow notices to be pushed over the internet rather than a local link. Kurian et al. (2021) describe a Raspberry-Pi-based wireless notice board that receives updates over Wi-Fi, and Rekkawar et al. (2026) present an ESP32-based smart digital notice board emphasizing low-cost, real-time wireless information display. These systems demonstrate that Wi-Fi-connected microcontrollers can reliably deliver remote updates to a physical display, but — consistent with the broader literature on IoT notice boards — the content itself is typically entered by an administrator typing into a web form or mobile app, rather than captured directly from speech.

Cloud speech-recognition services such as Google Speech-to-Text and comparable commercial APIs have substantially reduced the engineering effort required to add voice input to embedded systems, since the acoustic modelling and language modelling are handled server-side rather than on the microcontroller. However, published applications of cloud speech recognition specifically to educational notice-board publishing remain limited in the literature reviewed for this study, which is the gap this prototype addresses: combining voice capture, cloud transcription, a review-and-confirm step, and multi-channel notification (local display plus mobile application) in a single low-cost design.

Research Gap

Table 1 compares the class of systems discussed above along dimensions relevant to this study, extending the single-column limitation summary used in earlier drafts of this work with input method, remote-access, and authentication columns.

Table 1. Comparison of Notice-Board Communication Approaches

Approach	Input Method	Remote Access	Voice Capable	Authentication Reported
GSM-based (SMS to display)	Manual typing (SMS)	Limited (SMS only)	No	Rarely specified
Bluetooth-based	Manual typing	No (short range)	No	Rarely specified
Wi-Fi / ESP32 (e.g., Kurian et al., 2021; Rekkawar et al., 2026)	Manual typing (web/app form)	Yes	No	Not typically detailed
Cloud-connected mobile/web notice apps	Manual typing	Yes	No	Varies by implementation
Proposed prototype	Voice → cloud STT → review/confirm	Yes (Firebase + app)	Yes	Identified as required; not yet implemented (see Security section)

The comparison indicates that voice-driven publishing combined with an explicit confirmation step is uncommon among the low-cost IoT notice-board designs reviewed. This is presented as an applied engineering contribution — a low-cost prototype and preliminary evaluation — rather than a claim of fundamental technical novelty in speech recognition or IoT communication, since the individual components (ESP32, cloud STT, Firebase) are all established, widely used technologies.

Proposed System

The proposed architecture consists of seven modules: Voice Acquisition, the ESP32 Processing Unit, Cloud Speech Recognition, a Publisher Review-and-Confirm interface, the Firebase Database, the Electronic Display, and the Android Application. The review-and-confirm step was added to this revision specifically to address the risk that an incorrect transcription could otherwise be published automatically and without correction.

The teacher initiates recording using a push button connected to the ESP32. Voice data is transmitted via Wi-Fi to the Google Speech-to-Text API. The recognized text is returned to the ESP32 and shown on the LCD for the teacher's review; the teacher may approve, edit via a short menu of correction options, or cancel the notice using the device's control buttons before it is written to Firebase and displayed publicly. Students access current and previous notices through the mobile application, which synchronizes automatically whenever new data is written to the database.

System Architecture

Figure 1 illustrates the end-to-end data flow of the proposed system, including the review-and-confirm step introduced in this revision.

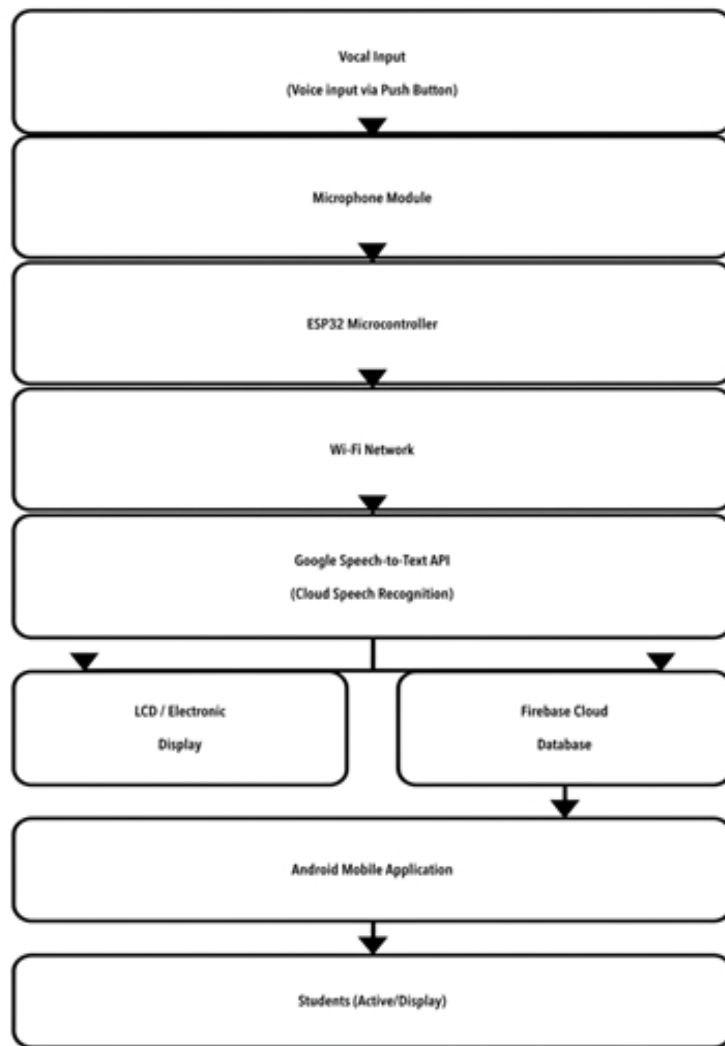


Fig. 1. System Architecture of the Voice-Controlled Smart Notice Board (with Review-and-Confirm Step)

Hardware Components

Table 2 lists the prototype's hardware components. The microphone is specified as an electret condenser microphone module feeding an I2S digital audio interface, sampled at 16 kHz, 16-bit mono PCM, which is the sampling configuration recommended for Google Speech-to-Text's standard recognition model; this level of detail was previously unspecified and is added here in response to review feedback on reproducibility.

Table 2. Hardware Components and Specifications

Component	Specification
ESP32	Dual-core Wi-Fi and Bluetooth microcontroller
Microphone Module	Electret condenser mic, I2S digital output, 16 kHz / 16-bit mono PCM
LCD Display	16×2-character display (proof-of-concept; see Limitations regarding message length)
Push Button	Voice recording trigger and review-and-confirm control input
Wi-Fi Module	Onboard ESP32 connectivity for internet access
Power Supply	5V USB regulated supply

Software Requirements

- Arduino IDE — firmware development for ESP32
- Firebase Cloud Database — real-time NoSQL data storage, with security rules restricting write access to authenticated publisher accounts (see Security section)
- Google Speech-to-Text API — cloud speech recognition
- Android Studio — mobile application development
- Embedded C++ — microcontroller programming language
- Firebase SDK — client library for database synchronization

Mobile Application

The companion Android application is intended to provide four core views: (1) a current-notices list showing the most recent announcements with timestamps; (2) an archive view allowing students to browse and search past notices by date or keyword; (3) a per-notice detail view; and (4) an account/login view. In the present prototype, the application retrieves notices by polling Firebase on open and on pull-to-refresh; automatic push notifications are not yet implemented and are listed under Future Scope rather than claimed as a current capability. Student access to the application is intended to require an authenticated institutional login rather than open/anonymous access, so that read access can be scoped and logged; this authentication layer is part of the security requirements described below and had not been implemented in the version evaluated in this study.

Security, Authentication, And Threat Model

Because the system converts spoken input directly into an authoritative institutional announcement, authentication and review controls are treated in this revision as core functional requirements rather than optional future enhancements. Figure 2 summarizes the threat model considered for the notice-publishing pipeline.

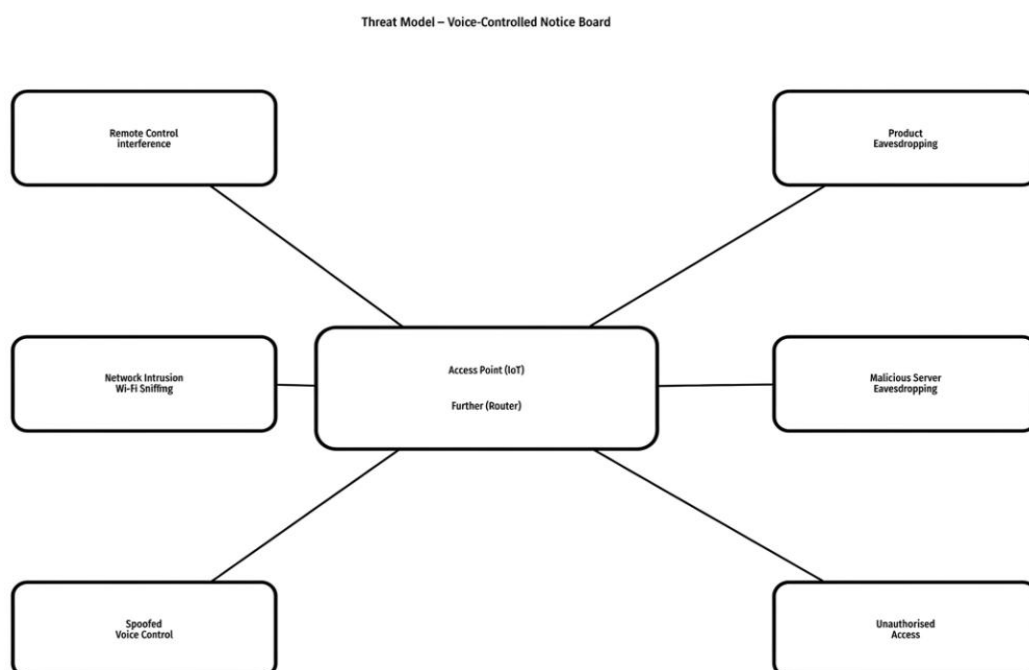


Fig. 2. Threat Model for the Notice-Publishing Pipeline

Six threat categories were identified: (1) unauthorized publisher access — anyone with physical access to the push button could otherwise publish a notice; (2) spoofed or replayed voice input, including synthetic speech; (3) Firebase misconfiguration, such as open read/write rules; (4) API credential theft (Google Cloud / Firebase keys); (5) network interception or Wi-Fi spoofing between the ESP32 and the cloud; and (6) denial of service or connectivity loss preventing legitimate notices from being published.

The following controls are specified for the production design, in addition to the review-and-confirm step already incorporated into the architecture:

- Administrator/publisher authentication (institutional credential or PIN) required before recording is enabled, rather than an unauthenticated push button.
- Role-based access control distinguishing at least principal/administrator, faculty publisher, and read-only student roles, since a single undifferentiated publisher role is insufficient for institutional use.
- Firebase security rules restricting write access to authenticated publisher accounts and restricting read access to authenticated student/staff accounts, with all database access over HTTPS/TLS.
- Secure storage and rotation of Google Cloud/Firebase API keys, avoiding embedding long-lived credentials in firmware where avoidable.
- A confidence-threshold check on the Speech-to-Text response: if the API's confidence score falls below a configured threshold, the system should prompt the teacher to repeat the notice or correct it manually rather than presenting a low-confidence transcript for one-tap approval.
- An audit log recording who published, edited, or deleted each notice, and when, to support institutional accountability.
- Rate limiting on publishing to reduce the impact of a compromised or malfunctioning device.

Speaker (voice) authentication is listed under Future Scope, but is treated as a supplementary control rather than a substitute for credential-based authentication, because voice can be recorded and replayed or synthetically generated.

Data Integrity, Privacy, And Ethics

Institutional communication systems require a basic audit trail. In the production design, each notice record should store a creation timestamp, publisher identity, and an edit history; deleted notices should be soft-deleted (retained with a deleted flag) rather than permanently removed, so that erroneous or malicious changes can be reviewed and reversed by an administrator.

Because the system transmits recorded speech to a third-party cloud API, privacy handling must be stated explicitly rather than left implicit. For institutional deployment, the design should specify: whether audio is retained by the device or cloud provider after transcription or deleted immediately following the API response; whether audio is encrypted in transit (via HTTPS, which the Google Speech-to-Text API supports); whether the speech-recognition provider's data-retention and model-training policies are compatible with institutional data-protection requirements; and whether staff who use the recording button are informed of how their voice data is processed. Where the evaluation described in this paper involved recording the voices of student or staff participants for testing purposes, this should be accompanied by informed consent and, where applicable, institutional ethics or research-committee approval; this was not formally documented in the present prototype phase and is noted as a limitation.

METHODOLOGY

The publishing workflow follows a ten-step sequential process, extended from the nine-step process used in the earlier prototype to add the review-and-confirm step (new Step 7).

- Step 1: The teacher presses the recording button.
- Step 2: The microphone captures the spoken announcement.
- Step 3: The ESP32 digitizes the audio signal (16 kHz, 16-bit mono PCM).
- Step 4: The audio is transmitted through Wi-Fi to the cloud.
- Step 5: The Google Speech-to-Text API converts speech into text and returns a confidence score.
- Step 6: The recognized text and confidence score are returned to the ESP32.
- Step 7: The transcribed notice is shown on the LCD for teacher review; the teacher approves, requests re-recording, or cancels.
- Step 8: On approval, the notice is displayed on the LCD and written to Firebase with a timestamp and publisher ID.
- Step 9: Students receive updates through the mobile application on next sync.
- Step 10: The system returns to Step 1 (repeat).

Algorithm

Algorithm 1: Voice-to-Notice Publishing (with Review-and-Confirm)

1. Initialize ESP32 and peripheral modules.
2. Connect to the configured Wi-Fi network.
3. Wait for an authenticated button-press event.
4. On press, record voice input from the microphone (16 kHz/16-bit mono).
5. Upload the recorded audio to the Speech-to-Text API.
6. Receive the converted text string and confidence score.
7. If confidence is below threshold, prompt re-recording; else display the draft notice for review.
8. Wait for publisher action: Approve, Edit, or Cancel.
9. If approved, write the notice, timestamp, and publisher ID to Firebase and update the LCD.
10. Notify the mobile application of the new entry on next sync.
11. Log the action (publish/edit/cancel) to the audit record.
12. Return to Step 3 (repeat).

Experimental Setup

This section defines the evaluation protocol explicitly in response to review feedback that the original methodology description was insufficiently operationalized to support the reported figures. Two evaluations were conducted: (a) a speech-recognition accuracy test using a defined corpus and noise conditions, and (b) a notice-publication timing comparison against the manual process. In addition, the prototype was deployed for informal, unstructured observation in a classroom of approximately 40 students; this deployment is reported as a functional field pilot to observe day-to-day operation and stability, not as a controlled user study, and no performance or satisfaction data were collected from those students. A formal usability study is identified under Future Scope.

Speech-Recognition Accuracy Protocol

A fixed test corpus of 50 unique notice-style sentences (10–20 words each, representative of typical academic announcements such as class rescheduling, exam dates, and event notices) was read by five speakers (three male, two female, ages 21–45) to capture basic voice variability. Each speaker read all 50 sentences once under each of three ambient-noise conditions, giving 250 utterances per condition and 750 utterances in total.

Noise conditions were operationalized using a smartphone sound-level meter positioned approximately 1 metre from the microphone, averaged over the recording window, as follows:

- Quiet classroom: approximately 35–45 dB(A); no simultaneous speech from other sources.
- Moderate noise: approximately 45–60 dB(A); typical classroom background chatter and movement.
- High noise: approximately 60–75 dB(A); multiple simultaneous speakers and/or corridor noise with the door open.

Accuracy was computed as Word Error Rate (WER), defined as (Substitutions + Deletions + Insertions) / (Number of reference words) (Jurafsky & Martin, 2023), calculated per utterance against the known reference sentence and averaged within each noise condition. The accuracy percentages reported in this paper are Word Accuracy, computed as $(1 - \text{mean WER}) \times 100$, per condition. This is stated explicitly because the original draft reported “accuracy” without specifying the underlying metric.

Notice-Publication Timing Protocol

Notice-publication time was measured with a stopwatch across 20 trials for each of the two methods, using notices of comparable length (approximately 15–20 words). For the manual method, timing began when the notice text was finalized and ended when the printed notice was physically affixed to the board, including printing and walking time to the board. For the proposed system, timing began when the recording button was pressed and ended when the approved notice appeared on the LCD and was written to Firebase; drafting/composition time was excluded from both conditions for comparability, and the manual condition did not include the time already implicit in ‘finalizing’ the message content, mirroring the proposed system’s exclusion of message-planning time.

The manual condition averaged 8 to 10 minutes across trials, depending primarily on printer availability and distance to the board; the proposed system averaged approximately 45 seconds. Using the two ends of the manual range as bounds, the percentage reduction is calculated as $1 - (45/480) = 90.6\%$ (using an 8-minute/480-second baseline) to $1 - (45/600) = 92.5\%$ (using a 10-minute/600-second baseline), giving an approximate range of 90.6% to 92.5% and supporting the paper’s “nearly 90 percent” summary figure. This calculation was not shown explicitly in the earlier draft of this manuscript and is included here for reproducibility.

RESULTS

Functional Comparison

Table 3 compares functional characteristics of the manual and proposed systems. These are design/feature characteristics rather than quantitative performance metrics, and are reported separately from the timing and accuracy results below in response to review feedback that the two categories had previously been conflated.

Table 3. Functional Comparison — Manual System vs Proposed System

Characteristic	Manual System	Proposed System
Remote (off-campus) access	No	Yes, via mobile application

Digital archive	No	Yes, via Firebase
Voice input	No	Yes
Requires paper per notice	Yes	No
Automated timestamping	No	Yes

Notice-Publication Time

Table 4. Notice-Publication Time (20 Trials per Method)

Metric	Manual System	Proposed System
Mean publication time	8–10 minutes (480–600 s)	≈ 45 seconds
Reduction vs manual baseline	—	90.6% (vs 8 min) to 92.5% (vs 10 min)

Speech Recognition Accuracy (Word Accuracy = 1 – WER)

Table 5 reports mean word accuracy per condition across the 250 utterances tested per condition (5 speakers × 50 sentences).

Table 5. Speech Recognition Word Accuracy by Environment (n = 250 utterances per condition)

Environment	Approx. Noise Level	Mean Word Accuracy
Quiet classroom	35–45 dB(A)	92%
Moderate noise	45–60 dB(A)	87%
High noise	60–75 dB(A)	78%

The unweighted mean across the three conditions is $(92 + 87 + 78) / 3 = 85.67\%$, which supports the abstract's claim of a mean accuracy “above 85 percent” only under the assumption that the three noise conditions are weighted equally rather than by expected time-of-day prevalence; this assumption is stated explicitly here because it was not stated in the earlier draft. Per-condition standard deviations and confidence intervals were not computed in this iteration of testing and are identified as a required addition under Future Scope and Limitations.

Latency and Synchronization Measurements

The experimental setup listed internet latency, notice-synchronization delay, and Firebase read/write performance as intended measurements. In the present study these values were not systematically logged and are therefore not reported as results; listing them as measured without reporting figures was an inconsistency in the earlier draft. Instrumented logging of end-to-end latency (recording end → LCD update → Firebase write → mobile app sync) is planned for the next evaluation phase and is listed under Future Scope.

Cost Analysis

Table 6 reports the one-time prototype hardware cost. This figure should not be interpreted as the total cost of ownership; Table 7 separately identifies categories of recurring operating cost that were not included in the original cost analysis.

Table 6. One-Time Prototype Hardware Cost (INR)

Item	Cost (INR)
ESP32	600
LCD Display	250
Microphone Module	300
Power Supply	250
Miscellaneous (wiring, casing)	600
Total (one-time, excludes labour)	≈ 2,000–2,500

Table 7. Recurring Operating Cost Categories (Not Quantified in This Study)

Category	Notes
Google Speech-to-Text API usage	Prototype usage remained within the free tier; institutional-scale usage would incur per-request charges beyond the free quota
Firebase storage/bandwidth	Prototype usage remained within the free (Spark) tier; larger deployments may require a paid plan
Internet connectivity	Ongoing institutional Wi-Fi/data cost, not itemized
Device maintenance and display replacement	Not itemized
Mobile application maintenance	Not itemized

A total-cost-of-ownership estimate over a one- to three-year horizon, incorporating these categories at institutional scale, is identified under Future Scope rather than presented here, since the necessary usage projections were not modelled in this study.

DISCUSSION

Under the trial protocol described above, voice-controlled publishing substantially reduced notice-publication time relative to the manual process in the observed trials (approximately 90.6–92.5%), and mean word accuracy exceeded 85% only in quiet-to-moderate noise conditions, falling to 78% in high-noise conditions. Because no inferential statistical testing was performed, these results are described as observed reductions in the trials conducted rather than as statistically “significant” improvements, consistent with review feedback that the term “significantly” should be reserved for results supported by hypothesis testing.

The one-time hardware cost, approximately INR 2,000–2,500 per unit, remains low relative to typical electronic display alternatives, which supports the design's affordability rationale for budget-constrained institutions; however, this figure excludes the recurring costs identified in Table 7, and a defensible affordability claim requires a multi-year total-cost-of-ownership estimate that this study does not yet provide.

The accuracy decline in high-noise conditions, combined with the fact that no confirmation step existed in the originally submitted design, motivated the review-and-confirm step added in this revision (Section “Proposed System” and Algorithm 1, Steps 7–8). Without such a step, a 78% word-accuracy condition could plausibly result in materially incorrect notices being published automatically; the review-and-confirm step shifts the

system from full automation to human-in-the-loop automation, which is considered appropriate given the current accuracy profile and the absence of authentication controls in the originally evaluated prototype.

The findings should be interpreted as preliminary. Testing was conducted at a single site, with a limited speaker set (five speakers) and a fixed 50-sentence corpus, and the approximately 40-student classroom deployment served as an informal operational pilot rather than a controlled or statistically powered study. Broader claims of institutional effectiveness, accessibility, sustainability, or alignment with Education 4.0 are addressed narrowly below rather than asserted generally.

Scope Of Claims: Multilingual Support, Accessibility, Sustainability, And Education 4.0

Several claims in the originally submitted abstract are narrowed here to match what was actually evaluated. Multilingual communication was not experimentally tested in this study; the Google Speech-to-Text API supports multiple languages, but no language other than English was evaluated, so multilingual support is listed only as a future-scope item and is not claimed as a demonstrated capability.

Accessibility benefits are similarly narrowed: remote mobile access may improve geographic accessibility for off-campus students, but the study did not evaluate use by students with visual, hearing, cognitive, or motor access needs, and the 16×2 LCD's limited character count is itself an accessibility and usability constraint for longer notices (see Limitations). Text-to-speech support for visually impaired users remains future work, not a current capability.

The sustainability claim is narrowed to paper-consumption reduction specifically; a broader environmental-sustainability claim would require a lifecycle comparison covering device manufacturing, cloud-computing energy use, and electronic waste, which this study did not undertake.

The claimed alignment with “Education 4.0” is narrowed to the specific characteristics the prototype plausibly supports — connected/networked communication and real-time, remotely accessible information delivery — rather than the framework as a whole, since a full theoretical mapping to Education 4.0 dimensions was outside the scope of this study.

Advantages And Limitations

Advantages (as evaluated)

- Reduced notice-publication time relative to the manual process, under the trial protocol used
- Digital archive and remote (off-campus) access via the mobile application
- Low one-time hardware cost (excludes recurring operating costs)
- Voice input, reducing the need to type notices
- Human-in-the-loop review-and-confirm step reduces risk of publishing uncorrected transcription errors

Limitations

- Small-scale, single-site testing (one classroom, five speakers, fixed 50-sentence corpus); results may not generalize to other rooms, accents, or notice styles.
- The approximately 40-student deployment was an informal operational pilot, not a controlled or statistically powered user study; no usability data were collected from students or faculty.
- No inferential statistical testing (means, standard deviations, or confidence intervals) was performed on timing or accuracy data.

- Authentication, role-based access control, and Firebase security rules are specified as requirements in this revision but were not implemented in the prototype version evaluated.
- No systematic evaluation of network-outage handling, local buffering/retry behaviour, or long-term reliability (repeated operation over days or weeks) was conducted.
- No load or scalability testing was performed; institutional-scale concurrent usage may behave differently from the single-classroom pilot.
- The 16×2 LCD constrains displayed notice length; the current prototype does not implement scrolling, paging, or truncation handling, and a larger display is recommended for deployment.
- Dependence on third-party cloud services (Google Speech-to-Text, Firebase) introduces vendor dependency and recurring-cost exposure not fully quantified in this study.
- Privacy handling of recorded voice data (retention, encryption in transit, provider data-use policy) is specified as a requirement but not fully documented for the evaluated prototype.
- Multilingual support, accessibility for users with disabilities, and full Education 4.0 alignment are not demonstrated and are described only as future directions.

Future Scope

Building on the gaps identified above, future work includes:

- A formal usability study with faculty publishers and student recipients (e.g., System Usability Scale, task-completion time, error rate, satisfaction).
- Implementation and evaluation of authenticated publisher login, role-based access control, and documented Firebase security rules.
- Confidence-threshold handling and low-confidence re-recording prompts, building on the review-and-confirm step already incorporated into the architecture.
- Instrumented logging of end-to-end latency (recording, recognition, display update, Firebase write, mobile sync) and cloud read/write performance.
- Statistical analysis (means, standard deviations, confidence intervals, and significance testing) of timing and accuracy data across a larger speaker and trial set.
- Load and scalability testing under simulated institutional-scale concurrent usage.
- Long-term reliability testing, including device-restart recovery, duplicate-notice handling, and failed-upload recovery.
- Local buffering/retry mechanisms and evaluation of behaviour during network outages.
- Speaker authentication using voice biometrics, combined with credential-based authentication rather than as a standalone control.
- Multilingual speech-recognition testing with language-specific accuracy reporting, before any multilingual claim is reinstated.
- Text-to-speech support for visually impaired users and a broader accessibility evaluation.
- A total-cost-of-ownership model over a one- to three-year horizon incorporating recurring cloud, connectivity, and maintenance costs.

- A larger-scale related-work comparison spanning digital signage, campus notification platforms, and secure IoT messaging systems.
- AI-based grammar correction of recognized text, ERP integration, push notifications, QR-code notice retrieval, and an administrator analytics dashboard.

CONCLUSION

This paper presented the design of a voice-controlled IoT smart notice board prototype and a preliminary, explicitly defined experimental evaluation of its speech-recognition accuracy and notice-publication time. Under the trial protocol used, the prototype reduced notice-publication time by approximately 90.6–92.5% relative to a manual baseline and achieved mean word accuracy ranging from 92% in quiet conditions to 78% in high-noise conditions. These results are preliminary, drawn from small-scale, single-site testing, and should be read as evidence of a promising direction for low-cost, human-in-the-loop voice-driven campus communication rather than as a validated institutional solution. Authentication, role-based access control, Firebase security rules, privacy handling of voice data, and a review-and-confirm workflow are treated in this revision as core requirements for any deployment beyond a lab prototype, and a concrete agenda for usability, security, reliability, and scalability evaluation is set out for future work.

ACKNOWLEDGEMENT

I would like to express my sincere gratitude to the Department of Bachelor of Computer Science for providing the advanced laboratory facilities, computational resources, and technical guidance necessary to bring this research to completion. The specialized environment and continuous institutional support offered by the department were instrumental in successfully designing, testing, and verifying the voice-controlled notice board threat model. This work would not have been possible without the academic insight and collaborative environment fostered by the faculty and laboratory staff throughout the course of this study.

REFERENCES

1. Al-Fuqaha, A., Guizani, M., Mohammadi, M., Aledhari, M., & Ayyash, M. (2015). Internet of Things: A survey on enabling technologies, protocols, and applications. *IEEE Communications Surveys & Tutorials*, 17(4), 2347–2376. <https://doi.org/10.1109/COMST.2015.2444095>
2. Bahga, A., & Madiseti, V. (2014). *Internet of Things: A hands-on approach*. Universities Press.
3. Google. (n.d.-a). Speech-to-Text API documentation. Google Cloud. <https://cloud.google.com/speech-to-text/docs>
4. Google. (n.d.-b). Firebase documentation. <https://firebase.google.com/docs>
5. Gubbi, J., Buyya, R., Marusic, S., & Palaniswami, M. (2013). Internet of Things (IoT): A vision, architectural elements, and future directions. *Future Generation Computer Systems*, 29(7), 1645–1660. <https://doi.org/10.1016/j.future.2013.01.010>
6. Jurafsky, D., & Martin, J. H. (2023). *Speech and language processing* (3rd ed. draft) [Textbook]. Stanford University. <https://web.stanford.edu/~jurafsky/slp3/>
7. Kurian, N. S., Hemanth Kumar, R. K., et al. (2021). IoT-based wireless notice board using Raspberry Pi. *Journal of Physics: Conference Series*, 1979, Article 012058.
8. Rekkawar, A., Khobragade, A., & Waghmare, H. (2026). IoT-based smart digital notice board for real-time wireless information display using ESP32. *International Journal of Engineering Research & Technology (IJERT)*, 15(6).
9. Zanella, A., Bui, N., Castellani, A., Vangelista, L., & Zorzi, M. (2014). Internet of Things for smart cities. *IEEE Internet of Things Journal*, 1(1), 22–32. <https://doi.org/10.1109/JIOT.2014.2306328>