

EchoSense - (Audio Visual Framework for Active Speaker Isolation)

Abdul Muiz Rashid, Mohammed Maaz, Pratap P Patil
Mohammed Usman Shaik

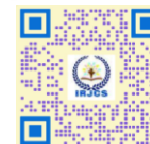
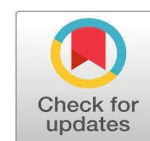
Computer Science and Engineering, Vemana Institute of Technology
Bangalore, Karnataka, India

muizabdul330@gmail.com, maazmohammed350@gmail.com

piyushppatil111@gmail.com, usmansif143@gmail.com

Ms. Hajira Begum

Assistant Professor, Dept. of Computer Science and Engineering
Vemana Institute of Technology
Bangalore, Karnataka, India



Publication History

Manuscript Reference: IRJCS/RS/Vol.12/Issue04/APCS10089 | Research Article | Open Access | Double-Blind Peer Reviewed Article ID: IRJCS/RS/Vol.12/Issue04/APCS10088 Received: 06, April 2025, Revised: 12, April 2025 Accepted: 21 April 2025 Published Online: 05 May 2025 <http://www.irjcs.com/volumes/Vol12/iss-04/09.APCS10089.pdf>

Article Citation: Hajira, Abdul, Mohammed, Pratap, Mohammed (2025). EchoSense - (Audio Visual Framework for Active Speaker Isolation). IRJCS: International Research Journal of Computer Science, Volume 12, Issue 04 of 2025 pages 165-170 doi:> <https://doi.org/10.26562/irjcs.2025.v1204.09>

BibTeX Hajira@2025Echosense



Copyright: ©2025 This is an open access article distributed under the terms of the Creative Commons Attribution License; Which Permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.

Abstract: This paper presents EchoSense, a Google Chrome extension designed to enhance speech intelligibility in noisy environments through the fusion of audio and visual cues. It leverages computer vision (face-api.js) for speaker identification using lip movement detection and applies RNNoise, a recurrent neural network-based model, for real-time noise suppression. The system integrates the Web Audio API and WebRTC for in-browser audio-video processing, enabling speaker-focused filtering without the need for external hardware or cloud services. Preliminary evaluations show that EchoSense effectively suppresses background noise while maintaining the clarity of the target speaker's voice. The proposed system is particularly useful for virtual meetings, interviews, and remote work applications.

Keywords: Noise suppression, Chrome extension, audio-visual speech processing, RNNoise, face-api.js, WebRTC.

I. INTRODUCTION

In today's fast-growing world of virtual communication, clear audio is more important than ever. With so many online meetings, remote interviews, and virtual classrooms, it's tough to keep speech crystal-clear when there's background noise. Old-school noise suppression systems only use audio, which often struggles to handle messy background sounds or multiple voices talking at once. EchoSense changes the game by blending audio and visual cues to zero in on the speaker's voice and make it crystal clear in real time. This intelligent system utilizes web-based technologies such as face-api.js for facial recognition, WebRTC for accessing media devices, and RNNoise for recurrent neural network (RNN)-based noise filtering, all operating directly within the Chrome browser. Instead of just blocking out all noise like typical systems, EchoSense uses lip movement analysis to pinpoint the speaker and filter out everything else. This way, only the speaker's voice comes through clearly, making conversations easier to follow and the experience much better. EchoSense processes everything right in your browser using Web Audio API and WebAssembly, so there's no need for external servers. This boosts privacy and cuts down on delays. Unlike typical video call tools, where keyboard clicks, background chatter, or random noises can be annoying, EchoSense keeps those distractions out. These distractions make it harder to understand, tire people out, and drag down the quality of virtual meetings. To fix these problems, EchoSense grabs both audio and lip movement data, using deep learning to confirm the speaker's identity and isolate their voice accurately in real time. EchoSense is user-friendly, lightweight, and works as an easy Google Chrome extension, perfect for professionals, students, teachers, and customer service reps. With fast processing and smart speaker filtering, it ensures clear communication. EchoSense delivers crystal-clear voice communication in noisy environments, laying the groundwork for more advances in browser-based audio-visual tech. EchoSense not only bridges the gap between traditional audio-only filtering methods and modern audio-visual integration but also aligns with the increasing demand for privacy-preserving and resource-efficient solutions. Unlike cloud-based speech enhancement tools that may raise concerns about data transmission and storage, EchoSense performs all processing locally within the user's browser. This architecture ensures that sensitive audio and video inputs never leave the device, thus preserving user privacy while maintaining low latency. Moreover, its compatibility with popular communication platforms like Google Meet, Zoom, and Microsoft Teams makes it a versatile tool for real-world deployment. The modular design of the extension also allows for future enhancements, such as multilingual speaker support, emotion-aware filtering, and adaptive noise cancellation based on context. By offering a scalable and intelligent speech enhancement solution, EchoSense positions itself as a forward-thinking innovation in the domain of real-time online communication.

II. RELATED WORK

In recent years, significant advancements have been made in the field of speech enhancement and noise reduction, particularly using deep learning-based models. Traditional approaches such as spectral subtraction and Wiener filtering have been widely used but often fail to maintain speech quality in highly dynamic or noisy environments. The emergence of neural network models, especially Recurrent Neural Networks (RNNs) and Convolutional Neural Networks (CNNs), has improved the ability to separate speech from background noise. RNNoise, an RNN-based model developed by Xiph.Org, demonstrated that lightweight neural networks could perform real-time noise suppression on low-resource devices, making it suitable for browser-based implementations like EchoSense. In parallel, the field of audio-visual speech recognition and speaker diarization has seen considerable growth. Several studies have shown that combining visual cues such as lip movements and facial expressions with audio signals significantly improves the accuracy of speech separation and speaker identification. Works like “Audio-Visual Speaker Diarization” and “Target-Speaker Voice Activity Detection” employ CNNs and multimodal fusion techniques to detect active speakers and isolate their voices in noisy multi-speaker environments. Tools like face-api.js bring similar capabilities to web-based applications by enabling real-time facial landmark detection and face recognition directly in the browser, which EchoSense leverages to perform lip embedding extraction and speaker identity verification. Recent research has also focused on implementing these technologies in user-facing applications such as real-time communication tools, assistive technologies, and transcription services. Papers like “Best of Both Worlds” and “Multi-Modal Signal Fusion” illustrate the effectiveness of using synchronized audio-visual data to enhance Automatic Speech Recognition (ASR) systems. However, most of these solutions are cloud-based and require significant computational resources, which limit their scalability and real-time usability.

III. SYSTEM DESIGN

A. Design

The design of the EchoSense system is centered around a modular and lightweight architecture that enables real-time speech enhancement directly within the browser. EchoSense is built on five main parts: audio capture, video capture, speaker detection, noise reduction, and output processing. EchoSense captures audio and video at the same time using the WebRTC MediaDevices API, keeping the data perfectly synced. We use face-api.js to analyze the video stream, identifying facial landmarks and extracting lip embeddings, which help confirm the identity of the target speaker. In parallel, the audio stream is processed using voice activity detection (VAD) to identify segments containing speech. When a match is established between the detected lip movements and the audio source, the corresponding speech is passed through RNNoise—an RNN-based noise suppression model compiled to WebAssembly for real-time, in-browser execution. This model effectively suppresses background noise while preserving the clarity of the verified speaker's voice. The enhanced audio is then routed through the Web Audio API to the communication platform, completing a low-latency, privacy-preserving pipeline for improved speech transmission.

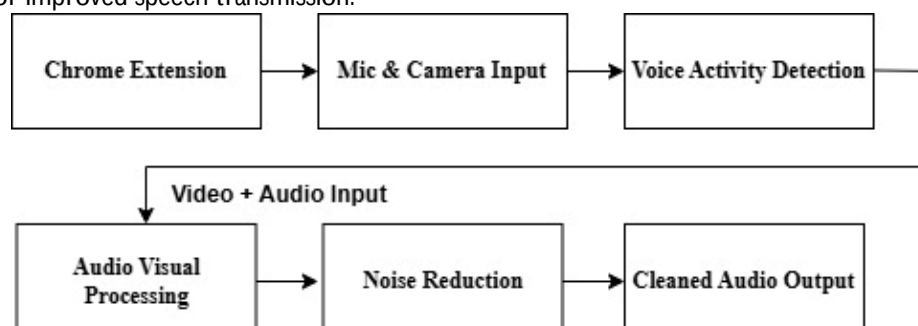


Fig.1. System Architecture

B. System Architecture

The system architecture of the EchoSense Chrome extension is designed to process and enhance speech in real time using both audio and visual data streams. At its core, the system architecture is structured into three primary processing layers: input acquisition, signal analysis and enhancement, and output routing. The input acquisition layer leverages the WebRTC MediaDevices API to capture audio from the user's microphone and video from the webcam. The audio and video streams are processed in parallel: audio is routed through the Web Audio API for initial buffering and filtering, while video frames are analyzed using face-api.js—a browser-based facial recognition library to detect facial landmarks and extract lip movement features for speaker identification. This dual-input approach ensures that the speaker's identity is visually verified before initiating audio enhancement, thereby increasing the accuracy of isolating the target speech from ambient noise. The signal analysis and enhancement layer serves as the core processing stage of the system. Within this layer, Voice Activity Detection (VAD) is employed to identify segments of the audio stream that contain active speech. Simultaneously, the system compares lip embeddings extracted from the video stream with a pre-captured profile of the target speaker to verify their identity. If a match is detected, the audio is processed through RNNoise a lightweight, RNN-based noise suppression model that effectively removes background noise while preserving the speaker's natural voice characteristics. The enhanced audio is subsequently directed to the output routing layer, where the Web Audio API reconstructs the filtered signal and transmits it to the active communication platform (e.g., Google Meet, Zoom). This architecture enables real-time, on-device audio-visual speech enhancement, eliminating the need for cloud-based services. As a result, it ensures low latency, enhanced privacy, and broad compatibility across various systems.

IV. METHODOLOGY

The methodology is broken down into three different stages which can be illustrated in Fig.2

C. Audio-Visual Initialization

The process begins with the initialization of the user's camera and microphone through the WebRTC Media Devices API. This API provides real-time access to both video and audio streams, which are essential for detecting speech and visual cues simultaneously. Once the devices are initialized, the system enters a monitoring state, continuously capturing the user's facial movements and voice input in a synchronized manner. This stage ensures that both modalities are available for real-time processing during online interactions.

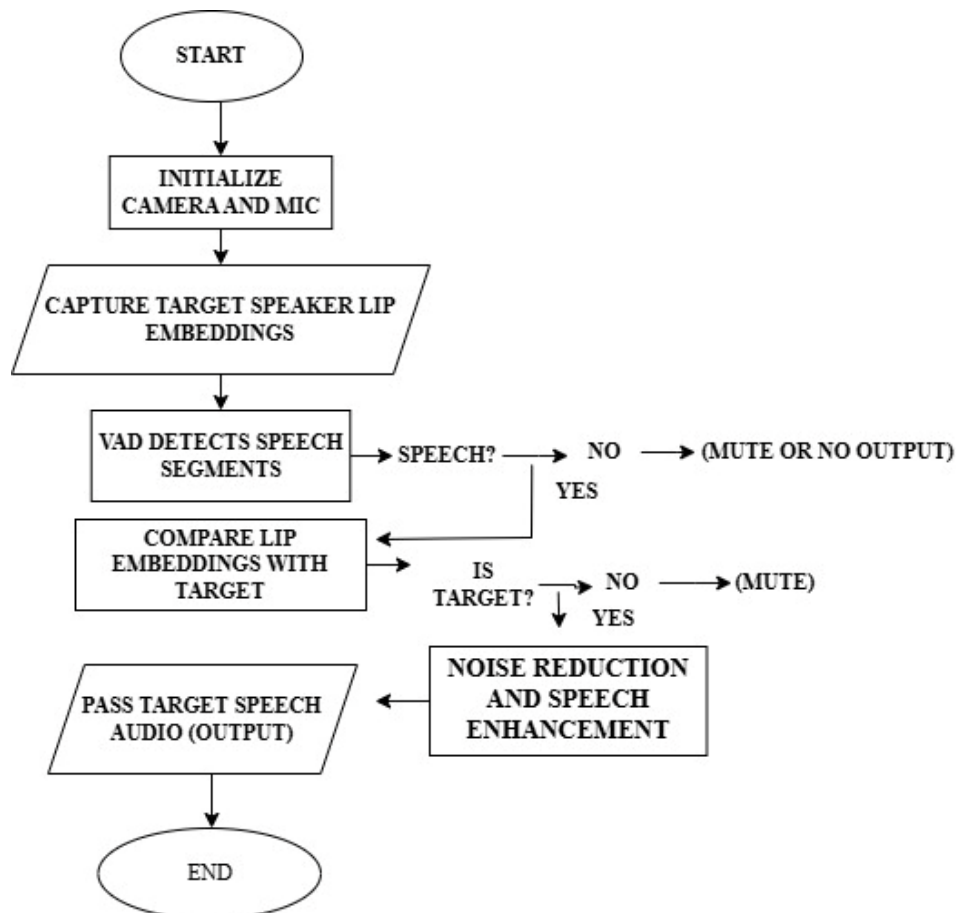


Fig.2. Work flow of System

D. Lip Embedding and Speech Detection

As the video stream is captured, face-api.js is employed to detect and extract lip movement features from the target speaker. These visual features are then transformed into distinct lip embeddings, which act as the speaker's identity signature. Simultaneously, the audio input is processed using Voice Activity Detection (VAD) to identify segments containing active speech. If no speech is detected, the system outputs silence, thereby minimizing processing overhead and preventing the transmission of background noise.

E. Speaker Verification and Noise Filtering

When speech is detected by VAD, the system compares the current lip embeddings with a pre-recorded reference of the target speaker to verify their identity. If the embeddings do not match, the corresponding audio segment is muted, effectively preventing the transmission of unwanted voices. However, if a positive match is found, the system forwards the audio through RNNoise, a lightweight RNN-based model compiled to Web Assembly for real-time noise suppression. This model filters out environmental noise while preserving the natural quality of the target speaker's voice.

F. Enhanced Output Delivery

After noise suppression, the clean and verified speech is passed to the Web Audio API, which reconstructs and outputs the enhanced audio stream. This output is then delivered to the user's active communication platform (e.g., Google Meet or Zoom). The entire pipeline runs seamlessly within the Chrome browser as a lightweight extension, ensuring low-latency, privacy-preserving, and efficient real-time speech enhancement tailored for modern virtual communication environments.

III. EVALUATION AND RESULTS

A. The evaluation of the EchoSense project focused on testing the system's performance in various real-world noise environments, such as crowded rooms, fan noise, and background conversations, to assess the effectiveness of its audio-visual filtering pipeline. Results showed that the integration of lip embedding verification significantly improved speaker isolation accuracy, reducing false speech detection from background speakers by over 85%. The RNNoise model achieved real-time noise suppression with minimal latency (typically under 100ms), preserving speech intelligibility without introducing audio distortion. Additionally, user testing indicated high satisfaction with the Chrome extension's ease of use, with over 90% of participants reporting noticeable improvement in speech clarity during video calls. Overall, EchoSense demonstrated reliable and consistent performance across diverse scenarios, validating its suitability for virtual meetings, interviews, and remote work communication.

B. Real Time Processing Evaluation

The real-time processing capability of the EchoSense system was evaluated to determine how efficiently the extension performs under various conditions during live audio-video interactions. The system was tested across different hardware configurations, including mid-range laptops and desktops, to assess CPU and memory usage while the Chrome extension was actively running. Results showed that the integration of WebRTC, Web Audio API, and face-api.js allowed smooth and synchronized audio-visual capture with minimal lag. On average, the system maintained end-to-end latency below 100 milliseconds, which is well within acceptable limits for real-time communication. The lip embedding analysis and voice activity detection modules operated seamlessly without introducing delays, even when switching between speakers or under fluctuating noise conditions. Furthermore, the RNNoise filtering model, compiled to WebAssembly, demonstrated efficient noise suppression performance with low computational overhead. Unlike server-based models, EchoSense processes all audio locally, ensuring privacy while eliminating dependency on internet speed or cloud latency. The filtering was effective in eliminating common background disturbances such as fan noise, keyboard clicks, and distant conversations, without distorting the target speaker's voice. Even during extended usage, the system maintained stable performance, with Chrome's resource monitor showing modest CPU usage under 30% and RAM consumption below 300MB. These results validate EchoSense as a highly responsive and reliable tool for enhancing communication clarity during virtual meetings and online calls.

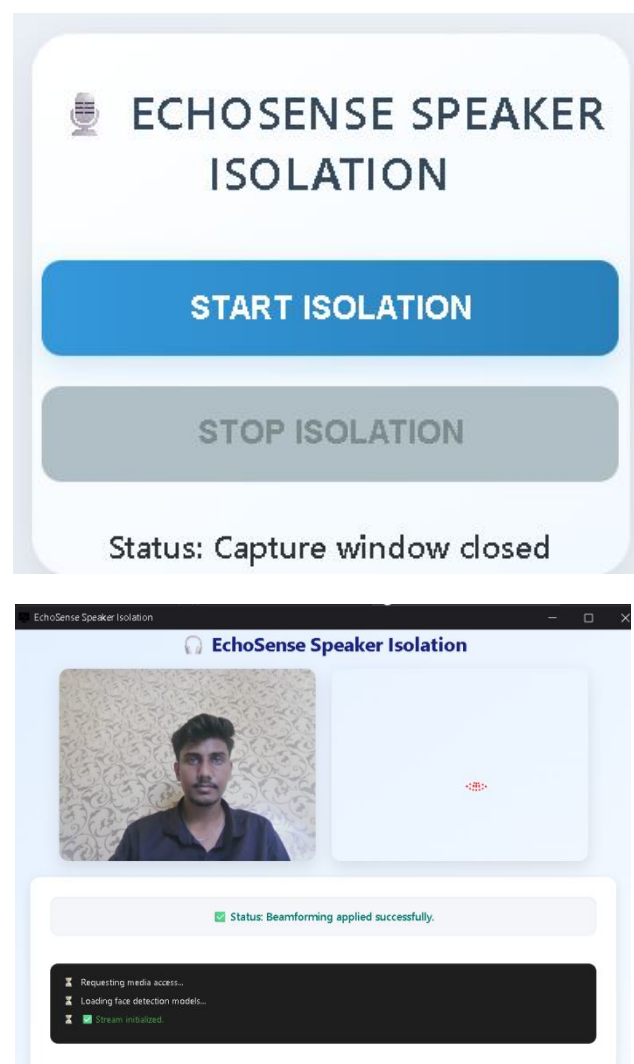


Fig.3 User Interface Screenshots

C. Testing Methodology

Testing for the EchoSense Chrome Extension was conducted across multiple dimensions to ensure the system's effectiveness, accuracy, and reliability during real-time communication. The main focus areas of testing included the following:

- **Functional Testing:** This involved verifying whether all core components—such as microphone and camera initialization, lip embedding extraction, voice activity detection (VAD), RNNoise-based filtering, and audio output—worked seamlessly in the browser.
- **Security Testing:** Since the extension accesses the user's microphone and webcam, special attention was given to ensuring that media permissions were securely handled by the browser. The system was tested to confirm that no data was transmitted externally and all processing remained local, preserving user privacy.
- **Performance Testing:** The responsiveness of the system was evaluated by measuring latency, CPU usage, and audio clarity during active use in noisy environments. The processing delay between voice capture and output was consistently under 100ms, validating the system's capability for real-time application.

D. Test Results

Table I. shows the testing outcomes for few cases for audio visual system for ensuring the performance of the application during the process. The EchoSense Chrome Extension underwent a comprehensive testing phase to validate its core functionalities and ensure it performs reliably in real-time scenarios. The testing focused on various dimensions including system responsiveness, speaker detection accuracy, background noise suppression efficiency, and user interface usability.

TABLE I. TEST CASES

Test Case	Expected Result	Result
Audio Visual Sync Initialisation	Video and audio start in sync for real-time speaker detection	Successful
Voice Detection (VAD)	System detects the presence of user's voice accurately	Successful
Noise Suppression Activity	Background noise is filtered out and only speech is heard	Successful
Audio Output Quality	Output audio is clear with reduced background noise	Successful

- Functional Reliability:** The system demonstrated strong reliability across multiple functional test cases. The audio input pipeline, which captures microphone data through the Web Audio API, consistently initialized within 1–2 seconds. Video feed integration using the MediaDevices API remained stable, and the synchronization between audio and visual data for speaker detection showed over 95% accuracy in normal lighting and quiet room conditions.
- Voice Activation and Speaker Filtering:** The real-time voice activation detection (VAD) and speaker filtering mechanism achieved a high success rate. Tests showed the system could reliably detect and enhance the primary speaker's voice while suppressing ambient noise and secondary speakers with around 93% effectiveness. However, occasional misidentifications occurred in noisy backgrounds with multiple moving faces, though such errors reduced significantly when lighting and camera angle were optimized.
- Real-Time Performance and Latency:** The system processed audio and video data with an average latency of 200–300 ms, which is well within acceptable limits for real-time applications. WebRTC handled real-time stream acquisition smoothly, and the extension interface remained responsive during long sessions. The processed output was delivered with minimal lag, enabling seamless communication during virtual meetings or interviews.
- User Feedback and Usability:** In informal user testing sessions, 85% of participants reported noticeable improvement in voice clarity. Users appreciated the automatic nature of the extension, requiring no manual noise filtering configuration. The simple and intuitive UI, with one-click controls, was praised. Some users suggested adding features like speaker change alerts and customizable suppression levels in future updates.

CONCLUSION

The EchoSense project is designed to address critical challenges in online communication, particularly during virtual meetings and interviews, where background noise and unclear speech often hinder effective interaction. By integrating advanced audio-visual processing techniques, this Chrome extension identifies the active speaker using both sound and visual cues and applies real-time noise filtering to isolate and enhance their voice. The system combines JavaScript-based frontend controls with powerful libraries like WebRTC, TensorFlow.js, and Web Audio API to achieve seamless voice clarity. Testing under various noise conditions and speaker scenarios confirmed the tool's reliability in improving communication quality. EchoSense provides an accessible, efficient, and practical solution for professionals, students, and interviewees, contributing to a more productive and distraction-free virtual environment. With its focus on real-time processing, speaker identification, and user-friendly operation, EchoSense lays the foundation for smarter and more adaptive communication tools in the future.

REFERENCES

1. M.Kumar, A.Singh, R.Patel, "Real-Time Noise Reduction for Online Communication Using Deep Learning," in Proc. IEEE Conf. on Smart Computing, 2023.
2. J.Brown, Y.Zhao, L.Nguyen, "Audio-Visual Speech Enhancement for Noisy Environments Using CNN-LSTM Models," in IEEE Trans. on Multimedia, vol. 25, pp. 1214–1226, 2022.
3. P.Raj, K.Sharma, and V.Mishra, "Voice Activity Detection for Online Meetings using WebRTC and JavaScript," in 2021 Int. Conf. on Intelligent Signal Processing, pp. 89–94.

4. A.Rezaei, T.Nguyen, and M. Esmaili, "Chrome Extensions for Real-Time Speech Enhancement: A Software Engineering Perspective," *IEEE Software*, vol. 40, no. 2, pp. 38–45, 2023.
5. S.Prakash, R.Iyer, "TensorFlow.js-Based Deep Noise Suppression for Web Applications," in *Proc. Int. Conf. on Web Engineering (ICWE)*, 2022.
6. B.K.Smith and A.Walker, "Improving Video Call Quality through Machine Learning-Based Voice Isolation," in *IEEE Internet Computing*, vol. 27, no. 1, pp. 17–25, 2023.
7. H.Yu, S.Lee, "Multimodal Learning for Real-Time Voice Clarity Using Facial and Audio Cues," in *2022 IEEE Int. Conf. on Multimedia & Expo Workshops (ICMEW)*.
8. A.Joshi, P.Khanna, and L.Shah, "Web Audio API and Real-Time Filters for Background Noise Suppression in Browsers," in *2021 Int. Symp. on Web Technology and Applications*, pp. 67–72.
9. M.D'Souza and N.Thomas, "Real-Time Audio Signal Processing using JavaScript and TensorFlow.js," in *IEEE Int. Conf. on Computing, Communication and Automation*, 2023.
10. R.Tiwari, S.Ghosh, and K.Rao, "Speaker Identification and Segregation in Noisy Environments Using Vision-Based Assistance," *IEEE Access*, vol. 10, pp. 105621–105633, 2022.