

Company Introduction | Product Introduction | Listen AI Performance | Use Cases

Listen AI

Sound-based emergency detection for urban safety monitoring

Deeply Inc. Introduction



Breakthrough Phase

(2025 - Present)

- 2025.03 Selected for HTX HATCH Dimension X in Singapore
- 2025.01 Registered as an Innovative Procurement Product
- 2025.01 Participated in CES and INTERSEC exhibitions



Growth Phase

2024 - 2023: B2B Sound AI

- 2024.11 Launched Open API
- 2023.02 Released emergency situation detection solution "Listen AI"
Secured over 10 business implementations and global customers



Foundation Phase

2017 - 2022: B2C Sound AI; R&D

- 2021.04 Released the parenting platform "BABBA"
- 2020.01 Launched mobile app "WAAH"
- 2017.07 Established Deeply Inc.

Major Awards

- 2024.07 Innovation of the Future Award (Korea)
- 2024.08 Incheon Proactive Administration Excellence Award (2024 H2)
- 2022 - Appointed to national data & public strategy committees
- 2021.12 ICT Innovation Square Pioneering Project Winner
- 2021.03 Member, Presidential 4th Industrial Revolution Committee
- 2019.12 Big Data Utilization Excellence Award (Korea Data Agency)
- 2018.12 Awardee, Gangwon Digital Healthcare Accelerator
- 2018.08 Startup Excellence Grand Prize (Jeong Juyeong CCEI)
- 2018.04 Winner, Digital Innovation Public Contest
- 2017.07 K-Global Startup Competition Grand Prize



Key Deployment Achievements

- 2025.03 **Hyundai Transys:** Abnormal fastening sound analysis
- 2024.09 **Korail:** Railway machinery anomaly detection
- 06 **Incheon Transit Corp.:** Chemical hazard detection
- 04 **Sejong Gov. Complex:** Cultural/sports facility hazard detection
- 2023.07 **DL E&C / Lotte E&C:** Construction site noise & hazard monitoring
- 05 **Incheon Startup Park:** Hazard monitoring service
- 04 **Kangwon Casino Land:** Resort hazard monitoring (High1)
- 2022.07 **Korea Army HQ:** Drone engine & anomaly sound detection
- 05 **Kookmin Bank, ITP:** Anti-fraud & anti-theft sound project
- 03 **KT (Korea Telecom):** Senior care & IoT communication projects
- 01 **Ministry of Health and Welfare:** 3,000 senior care devices deployed
- 2021.10 **LG CNS:** Emotion detection, military tech & sound AI projects
- 2020.09 **Amazon AWS:** Disaster response collaboration
- 2019.06 **Google:** Joint solution with Google for Startups
- 2018.09 **Meta:** Asia Startup Program (with Asia Foundation)



Proven Expertise in Sound AI Development

Intellectual Property Rights, Certifications, and Papers – Multiple Patents and Certifications

Intellectual Property Name	Country	Application / Registration Number
Real-time sound analysis method and system	South Korea	KR/10-2018-0075331 Patent No.10-2238307
Real-time sound analysis method and device	South Korea	KR/10-2018-0075332 Patent No.10-2155380
Machine learning method and device for automatic labeling	South Korea	KR/10-2018-0075333 Patent No.10-2189362
Data compression and encryption method and device	South Korea	KR/10-2018-0140185
Anomaly detection method	South Korea	KR/10-2018-0140186
Method, device, and system for distributed execution of deep learning models	South Korea	KR/10-2019-0129241
Efficient deep learning method and system for large-scale audio data	South Korea	KR/10-2019-0129242
Device, method, and program for spatially independent analysis of audio data	South Korea	KR/10-2020-0065199
Machine learning method and device for automatic labeling	PCT	OP19-0073PCT
Method and device for analyzing real-time sound	United States	US/36DEEP-BA10002PA
Baby sound analysis v1.0	South Korea	KR/BT-B-19-0167
Audio event analysis v1.0	South Korea	KR/BT-B-21-0080



- Real-time recognition of sound features and contextual analysis are core patented technologies
- Holds 3 patents in South Korea, 7 applications abroad, PCT patents, software certifications, and more.

Proven Expertise in Sound AI Development

Intellectual Property Rights

5



Globally Ranked No. 1 in Signal Processing Conferences (ICASSP*) with Numerous Published Papers



Google Scholar

Top publications

Categories > Engineering & Computer Science > Signal Processing >

	Publication	h5-index	h5-median
1.	IEEE Transactions on Image Processing	138	199
2.	IEEE Transactions on Wireless Communications	128	187
3.	IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)	123	198
4.	Conference of the International Speech Communication Association (INTERSPEECH)	107	162
5.	IEEE Transactions on Signal Processing	98	151
6.	IEEE Wireless Communications Letters	87	121
7.	IEEE Transactions on Circuits and Systems for Video Technology	88	120

UNPAIRED MICROPHONE CONVERSION FOR DEVICE GENERALIZATION IN SOUND EVENT CLASSIFICATION

*Myeonghoon Ryu¹ *Hongseok Oh^{1,2} Suji Lee¹ Han Park¹

¹ Deeply Inc.

² University of California, San Diego

ABSTRACT

In this study, we introduce a new augmentation technique for improving sound event classification (SEC) systems through the use of CycleGAN-based Microphone Conversion. We also present a unique dataset to evaluate this method. As SEC systems become increasingly common, it is crucial that they work well with audio from diverse recording devices. Our method tackles the issue of limited device diversity in training data by enabling unpaired training using spectrograms from various devices. The Microphone Conversion technology changes the style of the spectrogram while keeping the sound event details intact. Our experiments, conducted using 75 sound classes from 18 different devices, show that our approach outperforms existing methods in generalization by 5.2 - 11.5% in weighted F1 score. Additionally, it surpasses the current methods in adaptability across diverse recording devices by 6.5 - 12.7% in weighted F1 score.

Index Terms— Sound event classification, device mismatch, generative adversarial network, deep learning

1. INTRODUCTION

Sound event classification (SEC) aims to automatically identify various types of sounds, like speech, music, and environmental noise, using signal processing and machine learning techniques. Despite recent progress leading to practical applications like Apple's Sound Recognition and Alexa's Sound Detection, existing models still struggle with distortions caused by different recording devices[1, 2]. These distortions, although often subtle to human ears as shown in Figure 1, can significantly hamper SEC system performance[3]. Past strategies to tackle this domain shift from device variability have largely centered on data augmentation and normalization. Initiatives like the DCASE Challenge (Task 1)[3] have tackled this problem but have limitations due to their focus on synthetic evaluation data[4], reducing the thoroughness of its evaluations.

Previous studies have applied CycleGAN to tasks like speaker verification and emotion recognition[5, 6]. However,

*These authors contributed equally to this work.

they depended on datasets lacking detailed device information and did not use optimal recording environments like anechoic chambers. This oversight can introduce biases and challenges, especially when deploying on low-resource devices.

To bridge these gaps, we created a dataset featuring sound events recorded in an anechoic chamber using multiple real-world devices simultaneously. This dataset facilitates a deeper understanding of device-induced performance variations. Alongside, we have analyzed notable methods from successive DCASE Challenges (Task 1) and compared their performance against our approach, details of which are in Section 5.2.

Our key contribution is the introduction of Microphone Conversion, a plug-and-play augmentation technique for SEC systems. Utilizing the CycleGAN[7] framework, our method enables more effective training on multiple devices without requiring paired or labeled data. This plug-and-play solution generates spectrograms resembling recordings from desired devices. It empowers SEC systems to either generalize across devices or optimize for a specific one, showing enhanced performance compared to existing state-of-the-art methods.

2. RELATED WORK

3. DATASET

In this section, we outline the process of creating the Deeply Device Dataset and provide an overview of its contents. The dataset is generated by recording sound events using multiple devices concurrently within an anechoic chamber.

3.1. Sound Events

We have assembled a collection of 75 sound classes for our dataset. Of these, 25 sound classes were directly produced and recorded, while the remaining 50 are recordings of the playback of the "dry source" selected from the RWCP Sound Source Database (RWCP-SSD)[8]. For the former classes, each sound class lasts approximately 10 minutes. The RWCP-SSD contains non-speech sounds recorded in an anechoic room. These RWCP-SSD samples were down-sampled

"Microphone conversion: mitigating device variability in sound event classification"

: Advanced technology for improving sound classification performance across various devices (microphones).

World-leading (SOTA) algorithms for contextual analysis of sounds received from diverse microphones.

** Ranked No. 1 globally in signal processing (audio domain) conferences based on h-index criteria. Interspeech, ICASSP, WASPAA, DCASE

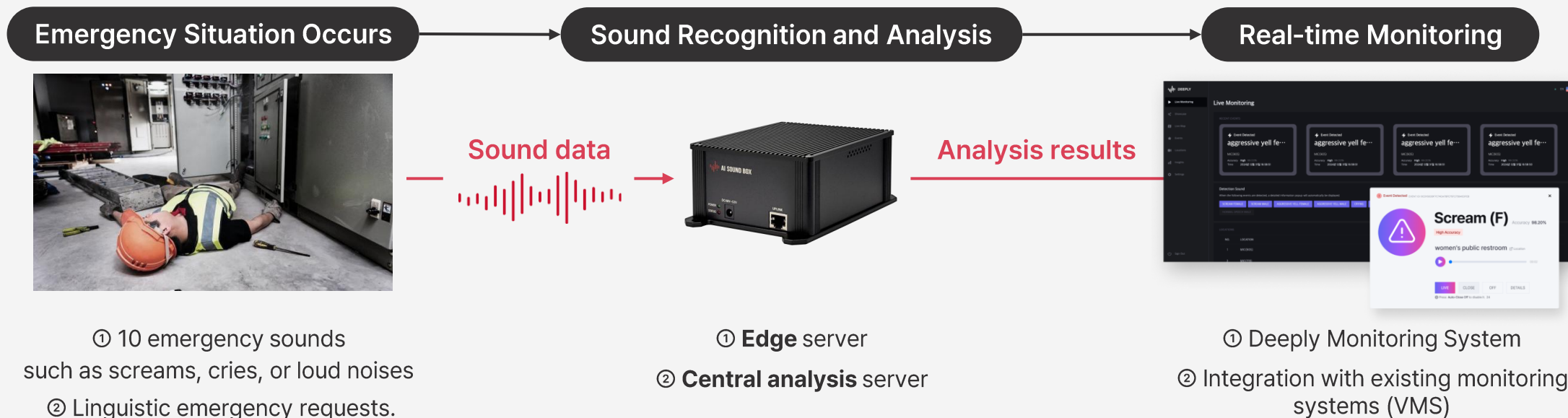
2024 Microphone conversion: mitigating device variability in sound event classification., Ryu et al.

2025 ViolinDiff : Enhancing Expressive Violin Synthesis with Pitch Bend Conditioning., Kim et al.

Noise-Agnostic Multitask Whisper Training for Reducing False Alarm Errors in Call-for-Help Detection., Ryu et al.

Sound-Based AI Emergency Situation Detection Solution

Listen AI



- ✓ **Continuous sound-based situation analysis for emergency detection.**
- ✓ **Integration with edge devices such as CCTV systems.**

Key Features of Listen AI

AI Monitoring Technology Utilizing Sound (Noise & Voice)



High Accuracy

- Detects real situational sounds—not just loud noises
- Industry-leading low false alarm rate
- Reliable detection even at a distance or in noisy environments (TV, music, etc.)



Exceptional User Convenience

- Real-time monitoring capabilities.
- Adjustable sensitivity for different sound types.
- Recognizes emergency linguistic requests only.

Key Features of Listen AI

AI Monitoring Technology Utilizing Sound (Noise & Voice)



Overcoming the Limitations of Video Surveillance

- Suitable for places with privacy concerns
- Reliable in any obstructions, lighting or weather condition



Solving Privacy Issues

- Transmits analysis results only, not raw audio (optional)



Advantages of Sound Surveillance

- Lower cost and easier setup than video systems
- Simple maintenance with high operational efficiency

Listen AI Sound Detection Event List

Real-Time Sound Detection Solution

9



Listen AI Safety v1.0

No.	Detected Sound Type	Details
1	Female Scream	Non-linguistic, emotional, and alarming scream
2	Male Scream	
3	Crying, Sobbing	Includes moderate to intense crying sounds
4	Breathing Heavily	Sounds of excessive breathing
5	Glass Break	The sound of large or small glass shattering
6	Physical Impacts	Sounds of objects falling or colliding
7	Female Normal Speech	Normal speech, focusing on conversational context
8	Male Normal Speech	
9	Female Aggressive Yell	Loud, aggressive yelling
10	Male Aggressive Yell	

Additional Situations and Sounds Detectable

(For inquiries, contact us)

Traffic-related Events
Accidents, speeding, etc

Speeding noise

Brake squeals

Collision sounds

Sirens

car horn



Smart Factory and Large-scale
Plant Monitoring for machinery
malfunctions

Bearing-related noise

Spindle

Fan

Specific abnormal sounds

Fastening noises

Air compressor

Intrusion Detection
Robbery, assault, fights, etc.

Heavy metallic movements

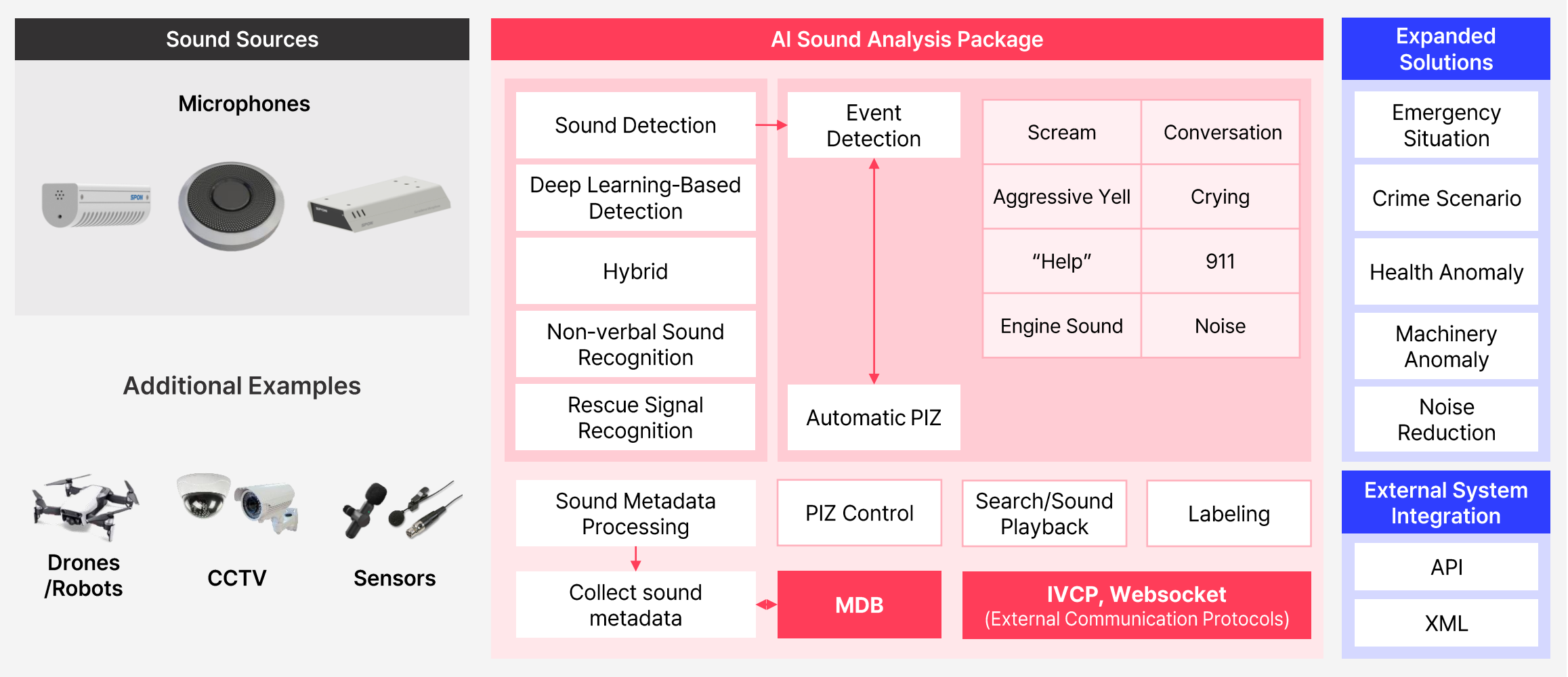
Specific phrases / distress calls

Multiple shot

Footsteps

Single shot





Ensuring AI Model Performance Through the World's Largest Nonverbal Dataset

Dataset 12



- Extensive dataset: 50,000+ hours of real-world audio
- Collected from diverse locations and acoustic conditions
- Designed for robustness across noise, distance, and devices

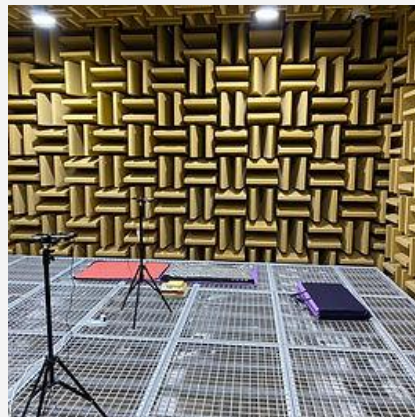
Data Collection Locations



<Restroom>



<Open Field>



<Anechoic chamber>

5 Key Large-Scale Dataset Features

Nonverbal audio

Emotional speech

Parent-child interaction audio

Various environments

Distance-based sound

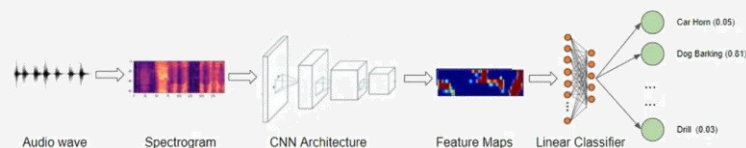
- Used to train and enhance ML/DL AI models
- Pre-trained models available per dataset

Listen AI

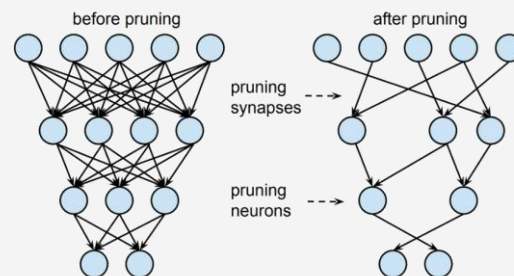
Artificial Intelligence

Deep Learning Technology

Not based on open-source models.
 Developed with **proprietary data** collected in-house.
 Creating **world-class AI models** using proprietary technology



- ✓ Development of **in-house AI models** optimized for various situations.



- ✓ Model optimization allows for both lightweight **CPU-based edge models** and **high-capacity server models**.

Strengths of Deeply's AI Models 13



- ✓ Patents in South Korea and the United States, academic papers, and **TTA-certified accuracy**.
- ✓ Achieved **98% F1-score**



- ✓ Ensures AI model performance with **optimized microphone processes**.

Key Features of Listen AI's Deep Learning Model Technology



Pre-trained Models

- High-performance, large-scale pre-trained models
- Optimized with 10+ Transformer-based architectures



Transfer Learning and Optimization

- Uses lightweight model frameworks
- Optimized for low-power and light-weight devices



Adaptability to Microphone, Space, and Noise Variations

- Trained with on-site field data
- Robust against environmental noise



Semi-supervised Learning

- Maintains stable performance in real-world settings
- Efficient even in compact environments



Product Introduction

Abnormal Sound and Voice Detection System



1. Edge Server

- Supports 4 channels.
- No data is transmitted to a central server, ensuring freedom from privacy issues

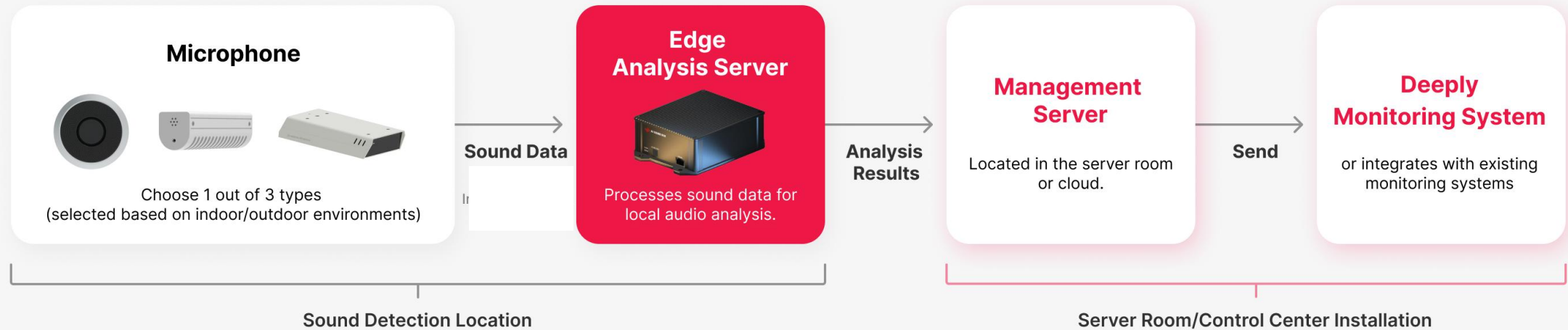


2. Central Server

- Supports over 100 channels.
- Operates by placing the server in the monitoring room



1. Edge Server



► Dedicated 3-Type Microphones

- High-performance microphones optimized for AI
- Three mic types for indoor, outdoor, and live use
- Powered and networked via Edge Server (PoE)
- PoE supports wired connections up to 70m

► Edge Server / Central Server

- Edge server analyzes sound locally to protect privacy
- Only analysis results are sent to the central server
- Supports up to 4 audio channels
- Mic-to-edge connection via PoE
- Requires separate power for edge server
- Edge-to-central networking via Ethernet

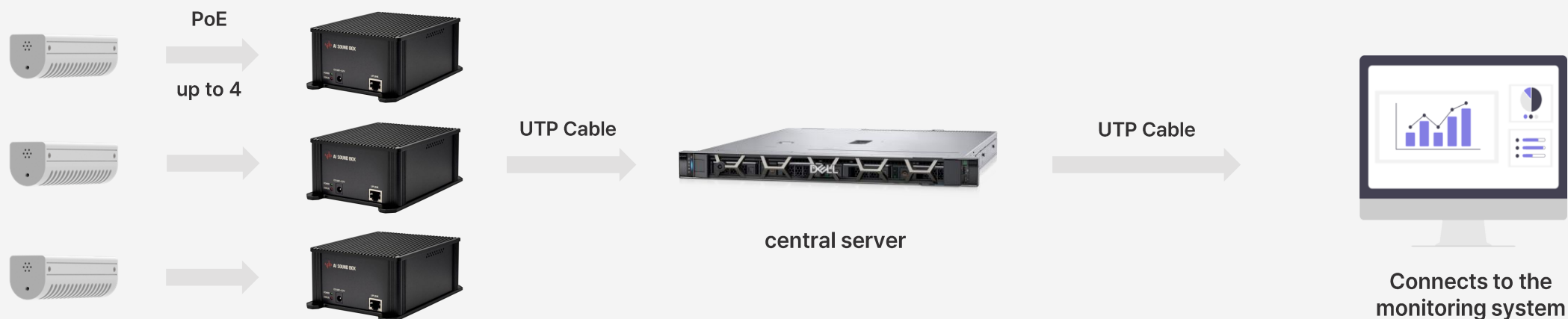
1. Edge Server : Choose between Standalone Edge Method and Multi-Edge Method



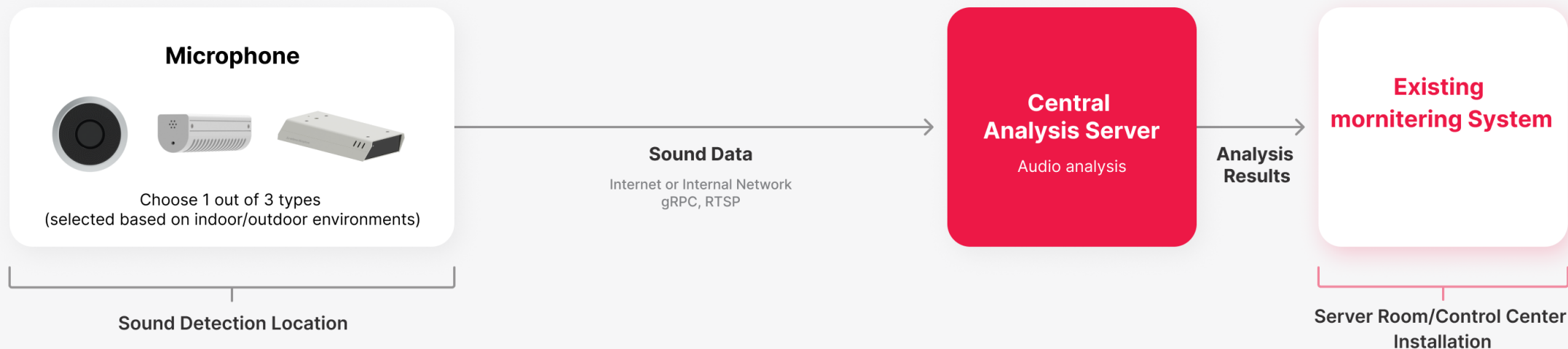
Standalone Edge Method



Multi-Edge Method



2. Central Server



► Dedicated 3-Type Microphones

- High-performance microphones exclusively for AI models.
- 3 options available based on indoor, outdoor, and live settings.
- Power and networking provided via PoE (Power over Ethernet).

► Central Server

- High-performance GPU analysis server
- Supports 50+ channel integration
- CPU: Quad-core 2.4GHz+ (x86_64)
- Memory: 16GB+
- Storage: 1TB+
- OS: Ubuntu 22.04 Server LTS or later

Hardware Configuration and Specifications

Dedicated 3-Type Microphones		
Image	Key Features	Specifications
	Indoor/Outdoor Use - Directional/Non-directional - Detection range up to 30m - Size: 11.3*4.2*3.7cm - Dynamic noise reduction - Compatible with CCTV	- 20Hz – 20KHz - Operating Temperature: - 20°C to 60°C - DC 12/1A, AC 24V, PoE (IEEE802.3at/af) - Power : <1W
	Indoor/Outdoor Use - Non-directional - Detection range: 5~180m² - Size: 8.5*1.8cm - Dynamic noise reduction	- 20Hz – 20KHz - SNR 90dB(@1KHz) - Operating Temperature: - 20°C to 60°C - Power Supply : DC 12/1A, PoE (IEEE802.3at/af) - Power : <1W
	Outdoor Use - Directional(optimized for long-range detection) - Detection range up to 20m - Size: 13.5*7.1*2.8cm - Dynamic noise reduction	- 20Hz – 20KHz - Sensitivity -35dB - SNR 90dB(@1KHz) - Operating Temperature: - 20°C to 60°C - IP54 weatherproof standard - Power Supply : DC 12/1A, AC 24V, PoE (IEEE802.3at/af) - Power : <1W

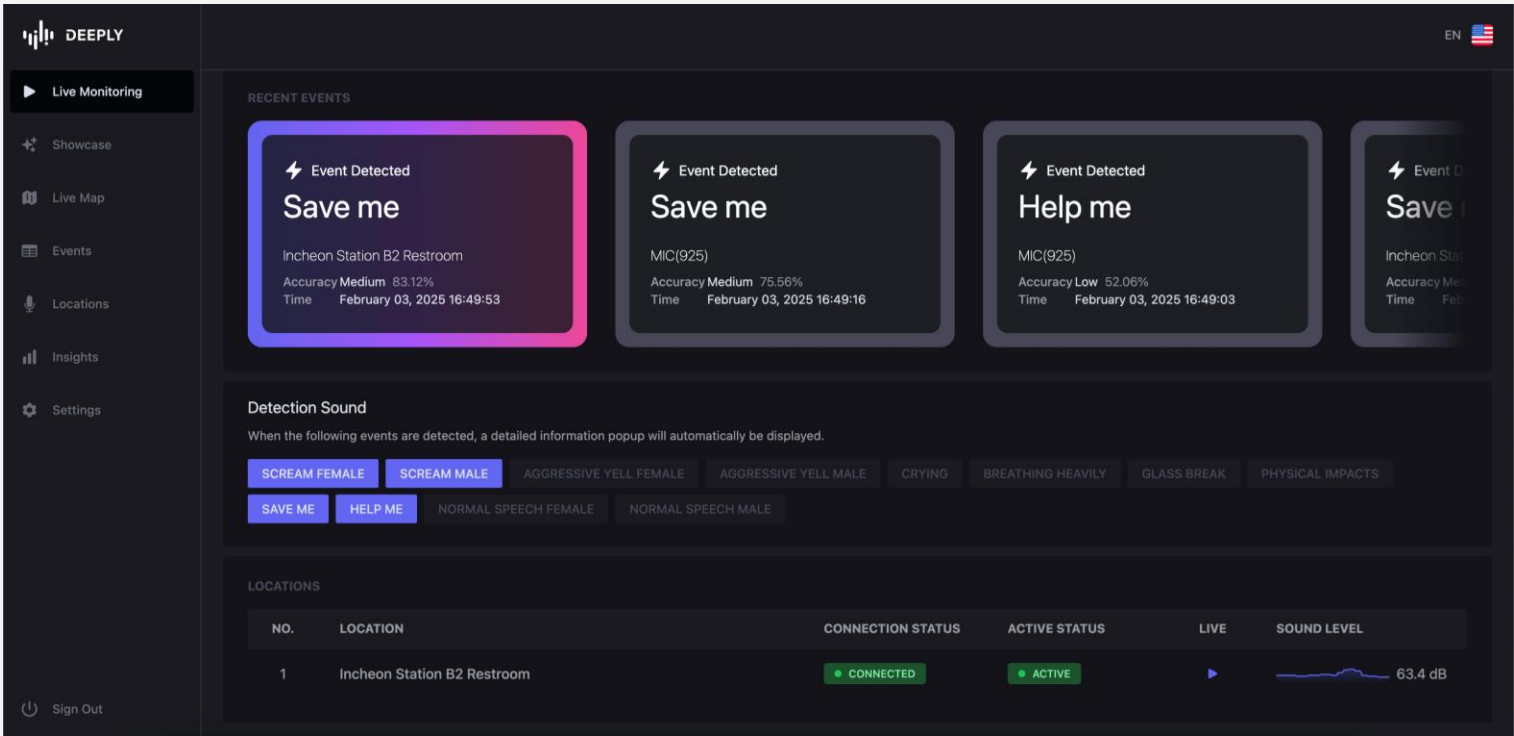
Server		
Image	Key Features	Specifications
	Edge server - Abnormal sound detection with AI. - Supports 4-channel microphone integration. - Size: 115×135×63mm	- Main Processor 4-core 2.4GHz + 4-core 1.8GHz or higher - RAM: 16GB - Storage: up to 32GB - Ethernet (PoE support)
	Verification server - Abnormal sound detection and AI engine operation. - Monitoring dashboard support	- CPU: 1CPU x Intel Xeon W-2425 3.00GHz (6-core/12-thread) processor - RAM: 64GB (32GB x2) DDR5 ECC-RDIMM 4800 Memory - Storage: 1TB SSD + 4 TB HDD - Power: Tower / 1200W x1 Fixed - Network: 1 x 10GbE / 1 x 1GbE - GPU: 2GPU x NVIDIA Quadro RTX A5000 - OS: Ubuntu
	Central Analysis server - Equipped with Abnormal Sound Detection AI Engine - Supports Monitoring Dashboard - Equipped with Abnormal Sound Detection AI Engine	- CPU: 1CPU x Intel Xeon E-2314 2.80GHz (4-core/4-thread) processor - RAM: 8GB (8GB x1) DDR4 - Storage: 1TB SSD - Network: 2 x 1GbE - OS: Ubuntu

Monitoring Dashboard

Integrated AI Analysis Dashboard

Providing a Dashboard for Real-Time Monitoring

- Classifies major sound types, including gendered voices
- Enables real-time risk detection beyond dB-based triggers



AI-Based Real-Time Emergency Detection and Dashboard

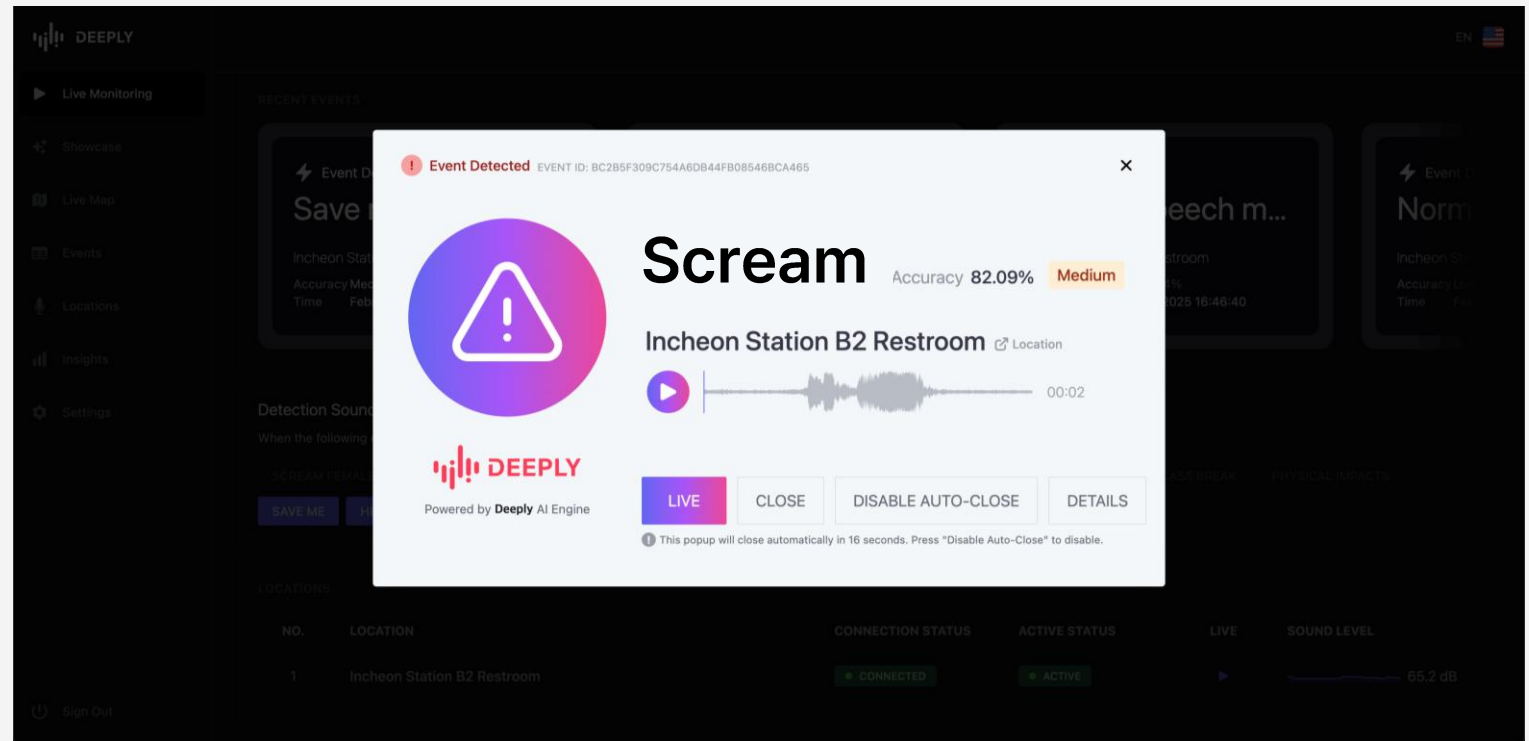
- Real-time display of AI-analyzed emergency events
- Shows mic status, connectivity, and sound intensity (dB)
- Web-based dashboard accessible from any browser
- Full system control: AI settings, alarms, and more

Monitoring Dashboard

Real-Time Event Display

Real-Time Emergency Detection and Event Popup

- When a suspicious sound related to an emergency situation occurs, a popup card is displayed on the screen.



Real-Time Event Popup Details

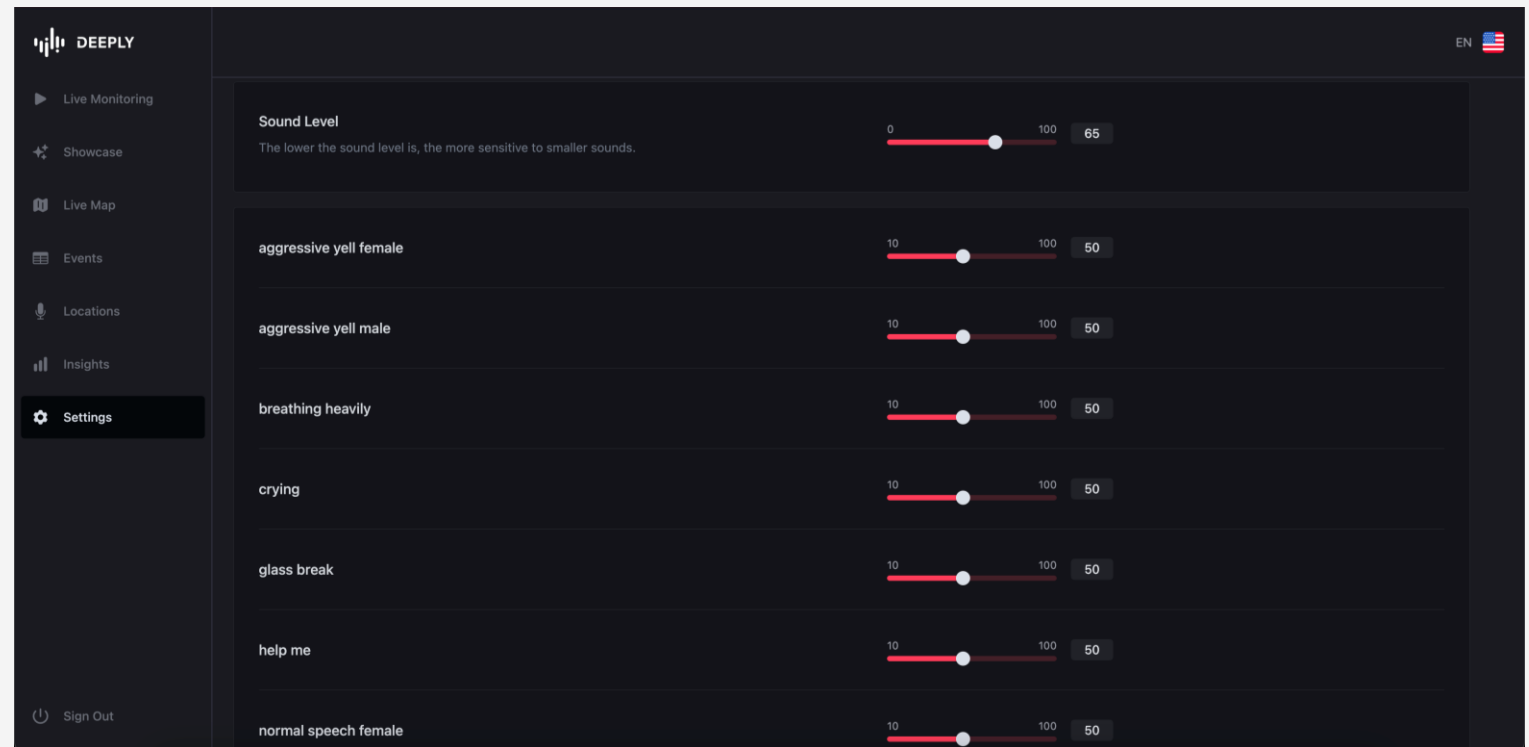
- Real-time emergency popups on the dashboard
- 2-second playback feature for quick situation assessment
- Instantly view sound type, location, and event details
- "More Info" button provides additional context

Monitoring Dashboard

AI Sound Detailed Settings

Optimized AI Analysis Settings for On-Site Conditions

- Fine-tuned detection by sound category
- Minimize false alerts in dynamic environments with smart threshold control.



Comprehensive AI Analysis Settings

- Configure minimum sound levels (dB) for sound event detection.
- Set confidence thresholds for each sound event type.
- Enable or disable location-specific analysis.
- Configure analysis times based on date, day of the week, or time of day.

Monitoring Dashboard

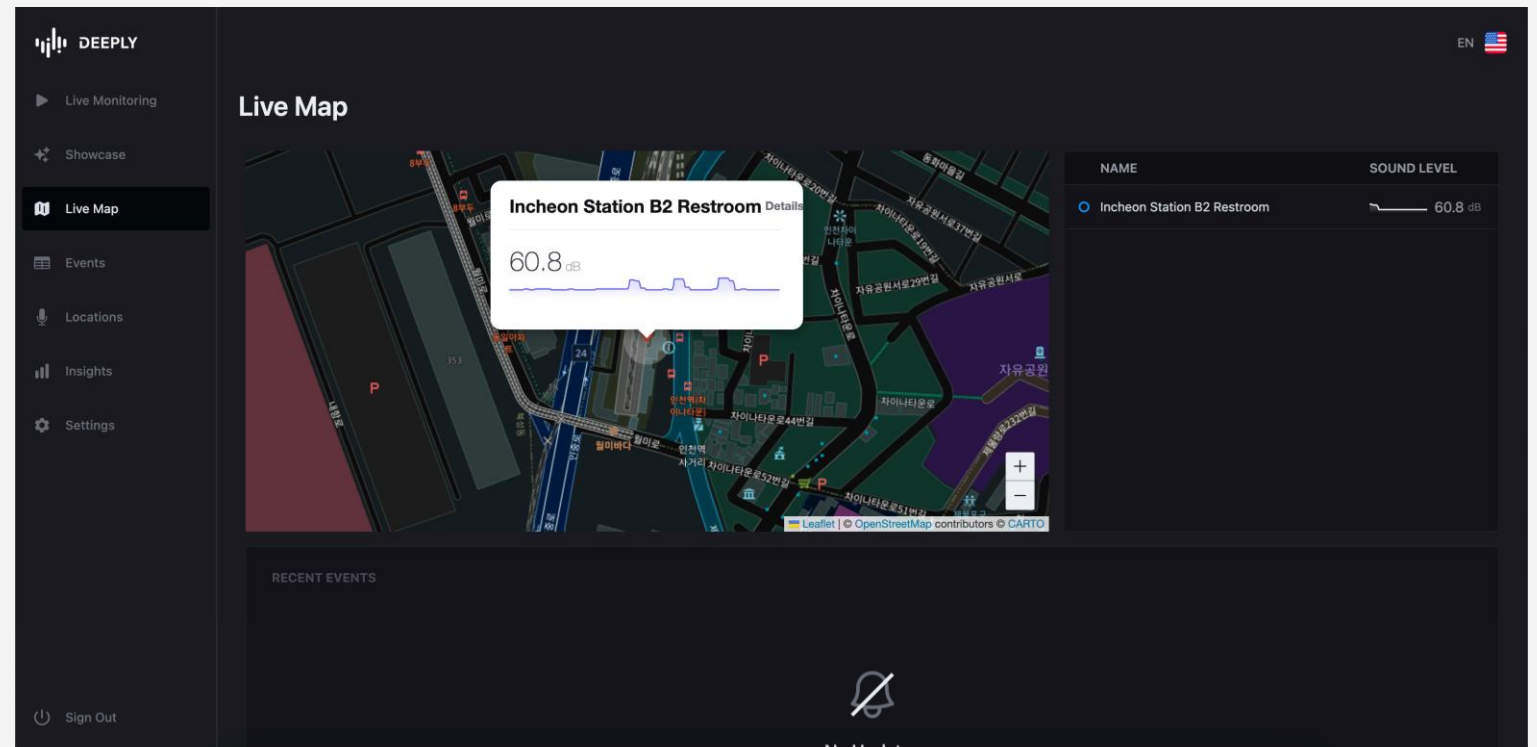
Location Verification and Live Streaming

Verify Installed Microphone Locations via Map

- Visualize mic locations on a real-world risk map
- Easily check multiple mic installations at a glance

Live Streaming (Overseas Support)

- Monitor microphone status in real time



Real-Time Microphone Verification Features

- Verify microphone locations on a map.
- Optional support for overseas locations.

Monitoring Dashboard

Key Features of the Dashboard

Events

EVENTS

Event

All

Location

All

Events per page

10

Date Range

NO.	TIME	EVENT	LOCATION
1	February 03, 2025 16:56:52	normal speech male	Incheon
2	February 03, 2025 16:56:51	normal speech male	Incheon
3	February 03, 2025 16:56:50	normal speech male	Incheon
4	February 03, 2025 16:56:49	normal speech male	Incheon

Event History

- Review events by date, device, and category
- View details: time, type, location, mic used, and confidence level

Insights

DAILY SUMMARY

Last

7 DAYS

30 DAYS

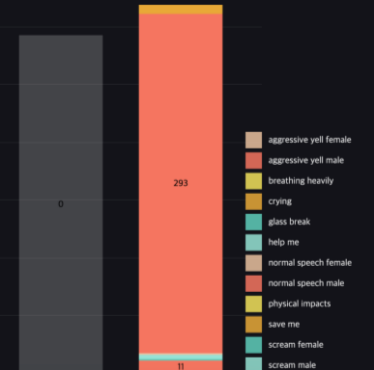
Event

AGGRESSIVE YELL FEMALE

AGGRESSIVE

NORMAL SPEECH MALE

PHYSICAL IMPAI



Statistical Analysis

- Offers stats on emergencies over selected time periods
- Shows trends and status by event type

Settings

PREFERENCES

Toggle Voice Alert

On

Voice Alert Settings

Voice Rate

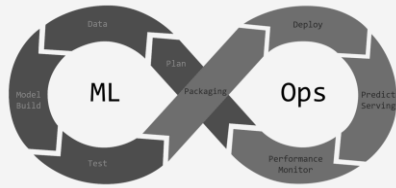
Voice Pitch

Play Test Voice

Settings

- Detailed configuration for AI emergency sound detection
- Customizable settings for microphones and other devices
- Includes license and system management tools

Achieving the Highest Accuracy and Lowest False Positives Among Current Technologies



MLOps Integration: Automated Learning System

- MLOps-enabled automated field data learning
- 10× accuracy improvement over standard models
- Robust and adaptable across diverse environments



High Accuracy, Low False Positives

- From 56,364 sound events in real-world cases, only 4 were false positives.
- Achieves **high accuracy** (F1 score of 98%) and minimal false positives.



Hardware Optimization

- Supports both CPU and NPU processing
- Optimized edge servers for 3 mic types and AI models
- Customizable mic setups for indoor/outdoor and space conditions

Product Design Prioritizing Privacy Protection

Observation Area



Sound Data



Physical Separation



Analysis Results



Control Center

Technical Safeguards for Privacy Protection

- Event Transfer Interface: Sends only detected sound types, not raw audio (optional)

Autonomous AI Operation

- On-device AI operates independently without external links
- Built to protect sensitive data from external exposure

Physical Network Segregation for Data Protection

- Dedicated interface ensures secure audio isolation
- Segregated mic connections for enhanced data protection
- Transmits insights only—raw sound never leaves the device



Open API

- Supports external integration via **Open API**.
<https://api-docs.deeplyinc.com>



- **ONVIF** protocol support (2025).
- Enables integration with standard security cameras, VMS, and other security/control programs.

SDK

- Embedded SDK for hardware integration (2025 release)
- Integrates Listen AI with devices like security cams and thermal sensors

** The installation method is provided in a separate document.



Listen AI Performance

Product Performance Test

: Ensuring performance in noisy environments

Recognition Rate Test by Noise Environment

(Recorded and tested in real-life indoor environments.)

3m radius			
	Noise	X	
		Water tap	Hand dryer
Scream	100%	100%	97.5%
Save me	100%	100%	90%

** Distress call detection has been tested in Korean

Accurate recognition
even in noise levels
up to 85dB

85dB is slightly louder than subway noise and is classified as "noisy operation" under industrial noise conditions.

Average accuracy in
noisy environments:

*rounded to the third decimal place

93.75%

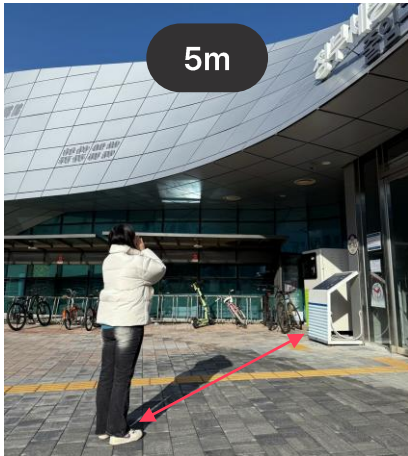
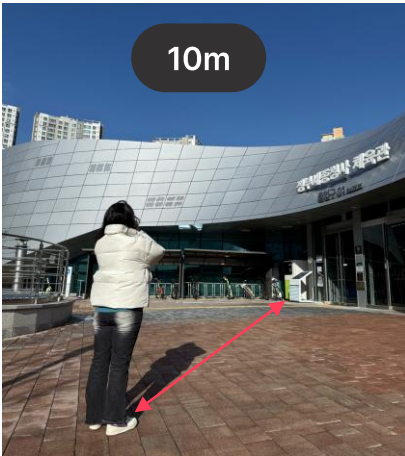
- Tested with consistent sound emission.
- Conducted by 4 male participants, each repeating the sound 10 times.

Product Performance Test

: Completed Sound Testing within a Radius of 20m

Recognition Accuracy Test by Distance

(Recorded in Outdoor Real-World Environment)

Outdoor				
	Scream	100%	100%	100%
	Save me	95%	95%	95%

** Distress call detection has been tested in Korean

Recognition Accuracy With No Errors Within Microphone Detection Range

*Test results based on TS-929E microphone with a detection radius of ~20m

Recognition Accuracy Difference by Distance 0%

- Experiments conducted with consistent sound intensity.
- Four adult males simulated emergency calls, repeated 10 times each.
- Standardized background noise: 55dB

Product Performance Test

: Minimized Errors through Deep Learning Technology

24-hour Indoor Error Test (Indoor Environmental Recording)



Office	*Analysis Attempts	Scream	Save me
All day	86,399	0	1
Working hours (9am – 7pm)	36,000	0	1
Otherwise	50,399	0	0
*noisy environment	1,799	0	0

24-Hour Analysis: Only 1 Error Detected Among 86,399 Events

Error Rate Within 24 Hours

*Among 86,399 detected events in 24 hours, only 1 error.

0.00001%

Error Rate During Office Hours

*Among 36,000 events detected during working hours (9 AM - 7 PM), only 1 error.

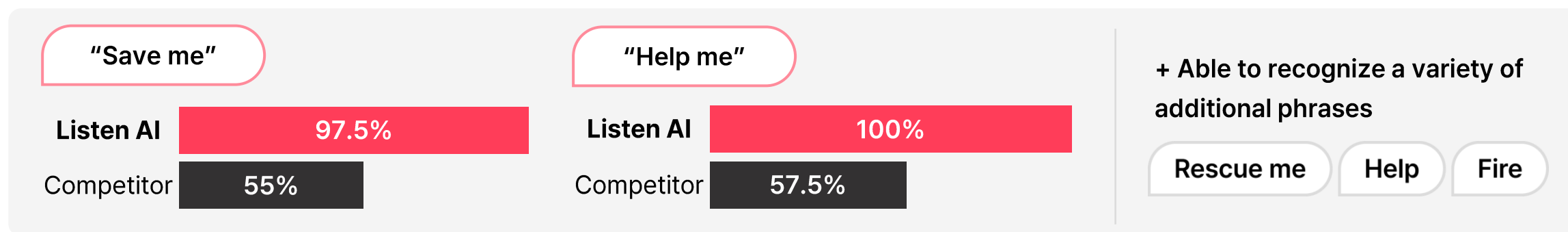
0.00003%

- *noisy environment
 - sustained noise levels above 80 dB for 1.5 hours.
 - 8 participants reproducing typical sound events
- *Analysis attempts
 - The number of analysis attempts made for detected events

Proven Listen AI: More Accurate and Versatile

Comparison of Recall (Detection Accuracy)

** Distress call detection has been tested in Korean



- Tested with 40 sound samples, including recorded and live female/male voices.
- Listen AI can even detect "Save me!" when spoken with a covered mouth, ensuring high accuracy in real emergency scenarios.

Detects **40 basic sound types** by default

Effective even for **muffled speech**

Completed testing in **diverse noisy environments**

"Confirmed superior performance compared to competing products."

Competitiveness and Excellence of Our Product

Verified Listen AI More Accurate, More Versatile

False Alarm

- Recorded During Tests Conducted 10 Times with the Same Sound

10x	Male Scream	Female Scream	Glass Break	Male Aggressive Yell	Female Aggressive Yell	Male Speech	Female Speech
Listen AI	0	0	0	0	0	0	0
Competitor	0	4	0	0	4	0	0

TTA, KOLAS Test Certification Secured

No.	시험항목	신청기관 기준	시험결과	No.	시험항목	신청기관 기준	시험결과
1	실내 사운드 기반 위급상황 감지 AI 모델 정확도	F1-score 90 % 이상	적합 (91.30 %)	1	키워드 검출 음성 AI 모델 정확도	F1-score 92 % 이상	적합 (92.29 %)
2	실외 사운드 기반 위급상황 감지 AI 모델 정확도	F1-score 85 % 이상	적합 (88.08 %)	2	성별 및 연령대 인식 AI 모델 정확도	F1-score 92 % 이상	적합 (93.28 %)
3	사운드 기반 위급상황 감지 AI API 초당 호출수	초당 호출수 RPS (Request Per Second) 1 RPS 이상	적합 (3.49 RPS)	3	키워드 검출 음성 AI 모델 API 초당 호출수	초당 호출수 RPS (Request PerSecond) 1 RPS 이상	적합 (3.52 RPS)



1. Accurate and Reliable

2. Proven Performance with TTA and KOLAS Certifications

3. Expandable with Additional Sound Sources and Training Capabilities



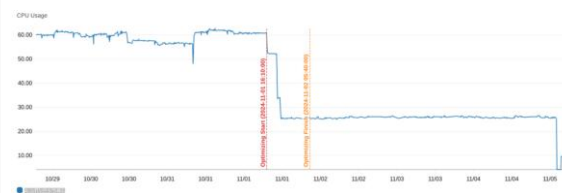
Use Cases

Sound-Based Risk Situation Monitoring Solution

- Installed in areas where CCTV installation is challenging, such as restrooms and emergency exits.
- Emergency monitoring for spaces frequently used by children and large crowds.

Model Advancement through Data Analysis

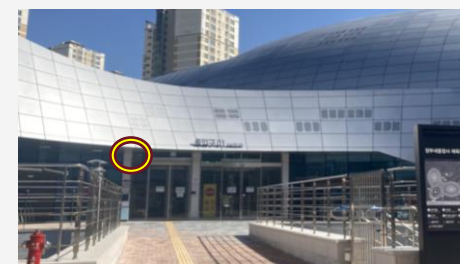
주차	날짜 구간	이벤트 수	비고
32	2024-08-11 ~ 2024-08-17	19,092	2024-08-14 설치
33	2024-08-18 ~ 2024-08-24	38,888	
34	2024-08-25 ~ 2024-08-31	37,113	
35	2024-09-01 ~ 2024-09-07	42,001	
36	2024-09-08 ~ 2024-09-14	44,456	
37	2024-09-15 ~ 2024-09-21	24,425	
38	2024-09-22 ~ 2024-09-28	56,364	
39	2024-09-29 ~ 2024-10-05	39,005	
40	2024-10-06 ~ 2024-10-12	20,020	
41	2024-10-13 ~ 2024-10-19	9,708	AI 모델 업데이트
42	2024-10-20 ~ 2024-10-26	12	연동 시작, 임계점 조절
43	2024-10-27 ~ 2024-11-02	5,604	내부 시스템 이슈 발생
44	2024-11-03 ~ 2024-11-09	4	연동 완료



- Applied AI model refinement to handle diverse noise factors (e.g., restroom ambient sounds, children's screams, break room air conditioner noises).
- On-site data integration enables model optimization for repeated and short audio events.

Installation Locations

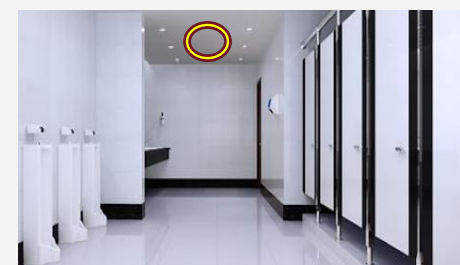
- Above Ground 1F** 3 locations - Main Entrance, North Restroom (Male/Female)
- Above Ground 3F** 7 locations - Restroom near Sauna, Emergency Exit (Male/Female), North Restroom (Male/Female), Break Room (Female)



Main Entrance



Changing room



Restroom



Sauna locker room



Real-Time Alerts for Emergencies in Subway Stations

- Increase in crimes and safety accidents in public restrooms due to blind spots in CCTV coverage.
- Existing CCTV-centric solutions struggle to detect all safety risks.
- Additional security systems are required to prevent crimes and monitor accidents in restrooms effectively.



Public Restroom Violent Crimes...

Increased 2.4 Times in Four Years



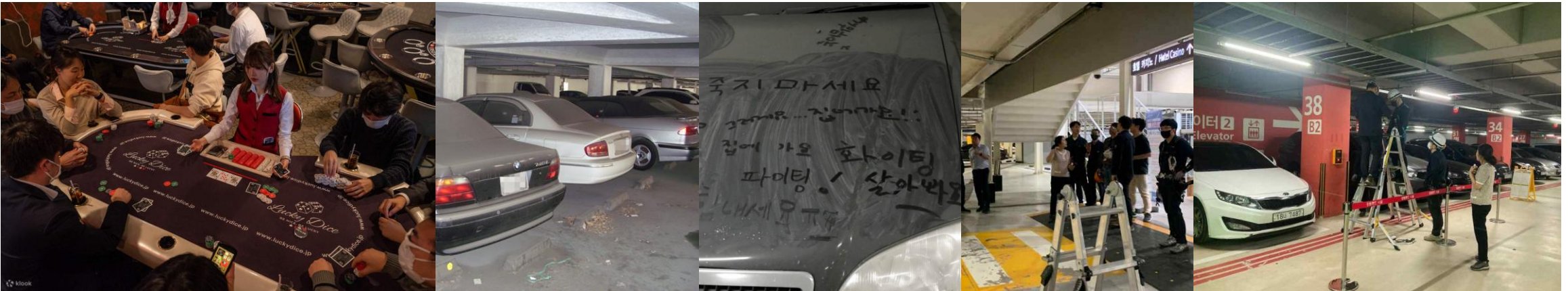
Sound-Based Risk Monitoring in Subway Station Restrooms

- Expanded sound-based risk monitoring projects in Bucheon and other areas planned for 2025.
- Quickly detects and responds to emergencies in subway stations to minimize damage and scale of incidents.

Real-Time Risk Sound Detection and Response Solution for Casino/Resort

: Responding to incidents such as crimes and loud noises in parking lots and smoking area

- Real-time detection of risk-related sounds in casinos and resorts
- Instant link to security teams for rapid incident response
- Improves operational efficiency and incident awareness
- Enhances customer safety and satisfaction



Implementation of Risk Monitoring System

- Updated AI detection model tailored for sounds such as screams and loud noises specific to Kangwon Land's needs.
- Generates monthly safety reports and provides alerts via a dedicated casino operation application.

Listen AI : Applicable Across Various Industries

Partners and Clients 39



Media Coverage Examples

Media Coverage 40



since 2022 deeplyinc.com 2024.07.24.

"Enhancing Sound Analysis with AI"...
Deeply Publishes Paper at
International Academic Conference

since 2022 deeplyinc.com 2024.09.20.

Subway Restroom Crimes and
Accidents Were a 'Nightmare'...
'Impressive Technology'

since 2022 deeplyinc.com 2024.07.24.

Deeply-Incheon Transit Corporation
Install 'Listen AI' in Restrooms at
Incheon National University Station

since 2022 deeplyinc.com 2024.07.24.

Deeply: Analyzing Sounds to Find
the Cause - "We Aim to
Contribute to the World"

since 2022 deeplyinc.com 2024.07.24.

CCTV Blind Spots... Ensuring
Citizen Safety Through Sound

since 2022 deeplyinc.com 2024.07.24.

"AI as a Safety Officer"... Real-time
Risk Detection and Prediction

since 2022 deeplyinc.com 2024.07.24.

"AI That Understands Sounds"
- Deeply Sets
a New Standard for Safety

since 2022 deeplyinc.com 2024.07.24.

"AI for Sound Analysis...
Limitless Applications"

since 2022 deeplyinc.com 2024.07.24.

Deeply Showcases Deep Learning-
Based Sound Analysis Safety Solutions
at 'Digital Innovation Festa 2024'

Selected for Singapore Hatch Dimension X Cohort 5

2025.02.20.

www.deeplyinc.com

Listen AI



Deeply, the second Korean company ever selected

Deeply has been selected as one of the 12 global companies from 7 countries including the U.S., U.K., Israel, working with Singapore's government agency HTX (Home Team Science & Technology Agency) on public safety PoC projects.

Deeply is emerging as a global leader in public safety tech through the Hatch X

Urban parks, gardens, and food courts are applying Listen AI's threat detection and monitoring technology.



Country	Solution
USA	Urban threat monitoring, gunshot and vandalism detection, tunnel intrusion detection
Türkiye / India	Urban threat monitoring, school safety monitoring
France	VIP estate safety monitoring
Mexico	Automotive factory machinery anomaly detection solution





Thank you

2nd Floor, 30 Baekbeom-ro 24-gil, Mapo-gu, Seoul, Deeply (04149)070-7459-0704

contact@deeplyinc.com

<https://www.deeplyinc.com/>