

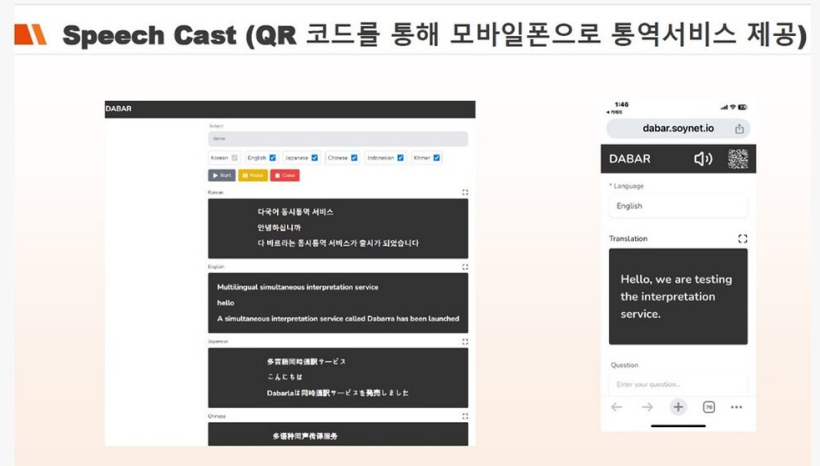
SoyNet Aims to Change Interpretation with Hybrid and On-Premise Solutions, with Private, Cost-Efficient Proprietary Tech

"Our Platform Reduces Data Leaks and API Costs by Deploying Lightweight, Secure Simultaneous Interpretation Systems for B2B, B2G, and Large-Scale MICE Events."

PANGYO, GYEONGGI-DO, SOUTH KOREA, August 20, 2026 /EINPresswire.com/ -- [SoyNet](#) Inc. is aggressively targeting the global MICE industry and the B2B and B2G interpretation markets, spearheaded by its AI-powered real-time multilingual simultaneous interpretation service, "[DABAR](#)". Established in 2018, SoyNet began as an AI inference accelerator developer and won recognition for its cost-reducing technical capabilities, leading to its selection for the prestigious TIPS program. Capitalizing on the rapid evolution and marketability of speech recognition technologies, the company conducted extensive research using translation engines such as Google Translate. Subsequently, it launched its cloud-based multilingual interpretation service, DABAR.



CEO Jung-Woo Park of SoyNet



DABAR's Speech Cast feature by SoyNet, which delivers interpretation services to mobile phones via QR code scanning | Image provided by SoyNet

DABAR is a comprehensive service that recognizes speech in 37 languages and delivers real-time translations in 124 languages. Unlike conventional interpretation apps that rely entirely on general-purpose APIs from global tech giants, SoyNet's DABAR implements its own foundational Speech-to-Text (STT), Machine Translation (MT), and Text-to-Speech (TTS) source

technologies—securing a competitive edge in data security and flexible pricing strategies. Notably, it overcomes long-standing hurdles in AI interpretation, such as rapid speech rates, individual pronunciation variations, and proper noun recognition errors, by deploying Large Language Models (LLMs) and personalized voice-learning functions to ensure contextually accurate outputs.

The service supports an integrated suite of features tailored for various environments, including “Speech Cast” for one-to-many broadcast interpretation, “Link” for one-on-one consultations, multilingual docents, mobile receivers, and multilingual meeting rooms. Furthermore, the company developed a dedicated hybrid interpretation platform for official events, proving its technical stability through successful operations at major international events, including GSAT in Changwon. This hybrid model features a collaborative framework where machine interpretation works alongside human interpreters and stenographers to correct interpretation errors in real time.

In addition, to fundamentally block corporate trade secrets or government confidential data from leaking through external APIs during the interpretation process, SoyNet has introduced an on-premise interpretation server that operates entirely within closed networks. This solution is gaining significant traction not only among tech firms and government ministries demanding top-tier security, but also among universities with large international student bodies that face steep financial burdens from per-account interpretation licensing costs. Religious organizations are also preparing to adopt this low-cost, high-efficiency infrastructure for multicultural ministries and overseas missionary operations.

Pangyo TV met with SoyNet CEO [Jung-Woo Park](#) at the Startup Campus in Pangyo Techno Valley to discuss their simultaneous interpretation ecosystem, on-premise security solutions, and upcoming global market expansion. Reporter Thomas Frederiksen conducted the interview.

Q1. Thomas Frederiksen, Reporter: Can you tell us a bit about your company? What kind of company is SoyNet?

A. Jung-Woo Park, CEO of SoyNet: SoyNet was founded in 2018. At the time of our launch, our artificial intelligence technology was about three times faster than Google’s inference speed, and we started our business with a core capability that could compress and lighten AI memory



SoyNet’s simultaneous interpretation feature deployed during a lecture conference at the Gyeongnam Science·Art·Technology Festival (GSAT 2026) | Photo provided by SoyNet

requirements down to one-fifteenth. We started by building AI inference accelerators that accelerate LLM training and reduce response latency. We possess this advanced technology for lightweighting and accelerating artificial intelligence models, and we are now pivoting that core engine into the interpretation space to solve foundational processing bottlenecks.

Q2. What is the meaning behind the name 'DABAR,' and how did you pivot your core optimization tech into the interpretation space?

A. Jung-Woo Park: The name 'DABAR' is derived from a Hebrew word that translates to "to speak" or "word." While looking for the optimal commercial domain in which to deploy our engine, we were concurrently researching voice recognition technology. As that voice engine fully matured, we determined that an audio-driven simultaneous interpretation service would be the absolute perfect application for our specialized acceleration pipeline. DABAR can recognize speech in 37 languages and output real-time translations in 124 languages. With this level of technical readiness, we re-architected and engineered a dedicated conference interpretation service based on real field lessons.

Q3. To successfully execute live convention interpretation, the AI must accurately track the rapid speaking pace of real-time presenters. How does DABAR solve these specific challenges?

A. Jung-Woo Park: A typical individual speaks at around 150 words per minute. However, professional announcers or conference presenters—especially when they are excited or nervous—frequently surge past 250 to 300 words per minute. When that happens, the translation task backlog rapidly grows. To address individual pronunciation discrepancies and speech rate variations, we invested heavily in R&D. For instance, our AI model applies contextual data smoothing to automatically filter and correct phonetic slip-ups in real time. Furthermore, we built a memory-injection feature that allows presenters to pre-upload their materials. The AI parses and understands the terminology in advance, enabling it to produce highly fluid, sophisticated translations during the live event.

Q4. Most translation tools look at single sentences, whereas a high-level conference environment requires deep context mapping. How is this multi-layered framework structured underneath?

A. Jung-Woo Park: Under the hood, DABAR operates through the synchronized orchestration of three separate AI engines. The pipeline is divided into three stages: first, converting live speech into text captions (STT); second, translating those captions into the target language (MT); and third, synthesizing those translated captions back into natural-sounding audio (TTS). To make real-time conversations possible, this entire multi-stage workflow must process within a fraction of a second. That rapid processing speed is where our core acceleration technology shines. Currently, we are continuously fine-tuning our engine to perfect the seamless audio-to-audio interpretation layer.

Q5. Corporate conferences and seminars routinely handle highly sensitive information. How does your on-premises interpretation server address data security constraints?

A. Jung-Woo Park: Corporate conferences and seminars routinely handle highly sensitive information, such as trade secrets, patent claims, or confidential government briefings. If a company relies on standard open-market APIs—such as Google, Microsoft, or DeepL—its translation data is transmitted globally to external servers where it can be used for machine learning and accidentally exposed to third parties. Realizing this critical vulnerability, we engineered a highly secure, completely enclosed ‘On-Premise’ installation interpretation server. Moving forward, delivering these secure, hardware-isolated translation servers to tech enterprises and security-conscious government bodies will be our core, front-line business focus.

Q6. I understand you are also extending this service to a completely different sector—the religious and immigrant community space. Can you share more about that?

A. Jung-Woo Park: I recently traveled to Australia and attended a local church service where they provided a QR code. When I scanned it, I received an instant, high-quality Korean translation of the sermon right on my phone. In South Korea, religious institutions tend to be quite conservative and hesitant to adopt AI tech, but the primary barrier is operational cost. To solve this, we engineered an incredibly cost-effective framework that utilizes lightweight infrastructure specifically optimized for religious environments. We are collaborating directly with major church denominations, and we project that around 1,000 churches will adopt our service in the near future. South Korea has rapidly transitioned into a multicultural society, hosting over 2.71 million migrants. Our solution aims to completely dismantle the language barrier.

Q7. Beyond the religious space, where else do you see this solution delivering the greatest economic value, and is your product already gaining market validation?

A. Jung-Woo Park: The highest value-add sector is definitely the MICE industry. Countries that form multi-ethnic hubs or host a massive influx of international business travelers—such as Singapore, Dubai, Saudi Arabia, Thailand, and Indonesia—present significant opportunities for us. We are already generating consistent, heavy revenue from our conference interpretation services. Our on-premises translation servers are gaining significant traction across diverse sectors, including local universities with large international student bodies, helping them reduce operational costs while maintaining full data privacy. Our technology has been rigorously field-tested and proven at South Korea’s top-tier international events, including the Global Startup Summit and G-START, as well as major regional platforms such as the Gyeongnam Science·Art·Technology Festival (GSAT 2026).

Q8. Finally, you are operating here in Pangyo Techno Valley. What are the core advantages of building your artificial intelligence company within this specific hub?

A. Jung-Woo Park: Pangyo offers an incredibly rich concentration of resources. The human capital and technical talent pool here are unmatched. Additionally, the localized ecosystem enables startups to access centralized government support and seamlessly engage with key technical acceleration programs, such as the global accelerator project supervised by the GBSA and operated by Y&ARCHER. Operating from the Pangyo Startup Campus, we are progressing with a funding round of approximately 1 billion KRW to scale our server capacity and sales force. Having access to prominent tech media platforms like AVING News provides an outstanding support engine for early-stage startups to secure vital global exposure.

Kim Seung Yeon

Gyeonggi Business & Science Accelerator

+82 31-776-4834

[email us here](#)

Visit us on social media:

[LinkedIn](#)

[Instagram](#)

[Facebook](#)

[YouTube](#)

[Other](#)

This press release can be viewed online at: <https://www.einpresswire.com/article/935451595>

EIN Presswire's priority is source transparency. We do not allow opaque clients, and our editors try to be careful about weeding out false and misleading content. As a user, if you see something we have missed, please do bring it to our attention. Your help is welcome. EIN Presswire, Everyone's Internet News Presswire™, tries to define some of the boundaries that are reasonable in today's world. Please see our Editorial Guidelines for more information.

© 1995-2026 Newsmatics Inc. All Right Reserved.