

SMEs Do Not Adopt Cyber Risk Management: Evaluating Voice-Based Conversational AI as a Barrier-Reduction Solution

MSc Research Project
Cyber Security

Cathal McGahan
Student ID: 23265884

School of Computing
National College of Ireland

Supervisor: Michael Pantridge

National College of Ireland
MSc Project Submission Sheet
School of Computing



Student Name: Cathal McGahan
Student ID: 23265884
Programme: Msc Cyber Security **Year:** 2025
Module: Msc Practicum/Dissertation/Thesis
Supervisor: Michael Pantridge
Submission Due Date: 11/12/2025, 2pm
Project Title: SMEs Do Not Adopt Cyber Risk Management: Evaluating Voice-Based Conversational AI as a Barrier-Reduction Solution
Word Count: 5911 **Page Count:** 19

I hereby certify that the information contained in this (my submission) is information pertaining to research I conducted for this project. All information other than my own contribution will be fully referenced and listed in the relevant bibliography section at the rear of the project.

ALL internet material must be referenced in the bibliography section. Students are required to use the Referencing Standard specified in the report template. To use other author's written or electronic work is illegal (plagiarism) and may result in disciplinary action.

Signature:

A handwritten signature in black ink, appearing to read "Cathal McGahan".

Date: 07/12/2025

PLEASE READ THE FOLLOWING INSTRUCTIONS AND CHECKLIST

Attach a completed copy of this sheet to each project (including multiple copies)	☐
Attach a Moodle submission receipt of the online project submission, to each project (including multiple copies).	☐
You must ensure that you retain a HARD COPY of the project, both for your own reference and in case a project is lost or mislaid. It is not sufficient to keep a copy on computer.	☐

Assignments that are submitted to the Programme Coordinator Office must be placed into the assignment box located outside the office.

Office Use Only	
Signature:	
Date:	
Penalty Applied (if applicable):	

AI Acknowledgement Supplement

Msc Cybersecurity Dissertation

Your Name/Student Number	Course	Date
Cathal McGahan (23265884)	MSc Cyber Security (MSCCYBE_JANOL24)	11/12/2025

This section is a supplement to the main assignment, to be used if AI was used in any capacity in the creation of your assignment; if you have queries about how to do this, please contact your lecturer. For an example of how to fill these sections out, please click [here](#).

AI Acknowledgment

This section acknowledges the AI tools that were utilized in the process of completing this assignment.

Tool Name	Brief Description	Link to tool
ChatGPT	Used quite extensively to help with the development of the POC, helping with Code and Prompt generation	ChatGPT (openai.com)

Description of AI Usage

This section provides a more detailed description of how the AI tools were used in the assignment. It includes information about the prompts given to the AI tool, the responses received, and how these responses were utilized or modified in the assignment. **One table should be used for each tool used.**

ChatGPT5.1	
Used quite extensively to help with the development of the POC code (React and n8n)	
I would be here for a while listing out all the prompts used, the majority of the prompts used were seeking development guidance, debugging and some explanations over groq, llms, n8n and other various pieces. AI was also used to help develop monologue scripting and test scenarios for the POC and help train the LLMs used using RAG and prompts. ChatGPT5.1 was used as a different baseline to test against, as detailed in the Evaluation section of research report. Listed below is a flavour of the types of prompts used to assist in the development and research of this project: • How do I set up n8n? • What is a n8n workflow? • How do I run n8n in a Docker container? • How do I set up Ollama locally? • How do I set up Groq? • What is RAG?	1 1. How do I set up n8n? n8n can be installed in several ways (desktop, CLI, Docker, cloud), but the most common method for local or self-hosted use is via Node.js or Docker. Option A: Install via Node.js (local development): 1. Install Node.js (LTS version recommended). 2. Install n8n globally: 3. npm install -g n8n 4. Start n8n: 5. n8n 6. Access the UI in your browser at: http://localhost:5678 Option B: Install via Docker (recommended for production): 1. Install Docker and Docker Compose. 2. Create a directory (e.g., /n8n/). 3. Add a docker-compose.yml file (example shown in next section). 4. Run: 5. docker compose up -d 6. Access the UI at http://localhost:5678. 2 2. What is an n8n workflow? An n8n workflow is a visual automation pipeline composed of nodes that perform specific operations. A workflow defines: • Trigger nodes (e.g., webhooks, schedules, events) • Action nodes (e.g., API calls, database operations, file transforms) • Data flows between nodes Workflows allow you to automate tasks, integrate systems, and

construct complex logic using a low-code interface.

3 3. How do I run n8n in a Docker container?

Below is a standard docker-compose.yml for running n8n:
version: "3"

services:

n8n:

image: n8nio/n8n:latest

ports:

- "5678:5678"

environment:

- N8N_BASIC_AUTH_ACTIVE=true

- N8N_BASIC_AUTH_USER=admin

- N8N_BASIC_AUTH_PASSWORD=yourpassword

- N8N_HOST=localhost

- N8N_PORT=5678

- N8N_PROTOCOL=http

volumes:

- n8n_storage:/home/node/.n8n

restart: unless-stopped

volumes:

n8n_storage:

Then start the container:

docker compose up -d

This will make n8n available at:
http://localhost:5678

4 4. How do I set up Ollama locally?

Ollama is a local LLM runtime that allows you to run models like Llama 3, Mistral, Gemma, etc.

Setup steps:

1. Go to <https://ollama.com/download>
2. Download and install the appropriate installer (macOS, Windows, Linux).
3. Verify installation:
4. ollama --version
5. Pull a model:
6. ollama pull llama3
7. Run a model interactively:
8. ollama run llama3

Optionally, Ollama can expose an API at http://localhost:11434 for integration with n8n, LangChain, etc.

5 5. How do I set up Groq?

Groq is a cloud-based LLM inference platform that provides extremely fast inference via their GroqLPU hardware.

Setup steps:

1. Create an account at <https://console.groq.com>
2. Generate an API key inside the dashboard.
3. Install the official Groq Python SDK (optional for dev work):
4. pip install groq
5. Use the SDK with your API key:
6. from groq import Groq
client = Groq(api_key="YOUR_API_KEY")
response = client.chat.completions.create(
model="llama3-70b-8192",
messages=[{"role": "user", "content": "Hello"}]
)
print(response.choices[0].message["content"])

Or use Groq via REST:

```
curl https://api.groq.com/openai/v1/chat/completions \  
-H "Authorization: Bearer YOUR_API_KEY" \  
-H "Content-Type: application/json" \  
-d '{  
  "model": "llama3-70b-8192",
```

	<pre>"messages": [{"role": "user", "content": "Hello"}] }'</pre> 6 6. What is RAG? RAG (Retrieval-Augmented Generation) is an LLM architecture pattern in which a model retrieves relevant external data (e.g., documents, PDFs, databases) and uses that retrieved context to produce more accurate outputs. RAG has two major steps: 1. Retrieve: Pull relevant information from a vector database or search index (e.g., Chroma, Pinecone, Weaviate). 2. Generate: Pass the retrieved context into the LLM so that responses are grounded in facts instead of relying only on the model's internal parameters. RAG is used for: • Chatbots over private documents • Knowledge-base question answering • Enterprise AI assistants • Compliance and audit workflows It improves accuracy, reduces hallucinations, and ensures outputs remain aligned with authoritative data.
--	---

Evidence of AI Usage

This section includes evidence of significant prompts and responses used or generated through the AI tool. It should provide a clear understanding of the extent to which the AI tool was used in the assignment. Evidence may be attached via screenshots or text.

Additional Evidence:

[Place evidence here]

Additional Evidence:

[Place evidence here]

SMEs Do Not Adopt Cyber Risk Management: Evaluating Voice-Based Conversational AI as a Barrier-Reduction Solution

Cathal McGahan
23265884

Abstract - SMEs experience cyber-attacks at high rates (43% annually) yet only 25% adopt risk management frameworks (n=16 survey). This research investigated whether voice-based conversational AI could reduce technological adoption barriers whilst maintaining acceptable accuracy within existing organizational constraints.

Mixed methods combined exploratory survey (n=16 SME owners, 7 industries) with technical evaluation of a voice-based system guiding users through CISRAM IGI assessment via natural dialogue. Survey validated five documented barriers (time, cost, complexity, expertise, awareness) whilst revealing non-adoption stems from competing priorities rather than awareness gaps. Technical evaluation compared AI-assisted completion against manual baseline across time, cost, and accuracy dimensions, testing two model sizes (20B vs 120B parameters).

Results demonstrated substantial barrier reduction. 24-fold time improvement (13 minutes vs 4-6 hours), 200-fold cost reduction (€1 vs €200-250), and eliminated expertise prerequisites. The 120B model achieved 87% accuracy with 6% mean absolute deviation, validating technical feasibility. Findings revealed system architecture determines performance more than computational scale (3-fold accuracy advantage over unstructured approaches), and that LLM evaluation requires multiple validation runs rather than single-test metrics. However, technological barriers must be reduced to thresholds enabling coexistence with competing priorities: 78% of non-adopters expressed willingness to adopt tools requiring under one hour, validating that barrier reduction to sufficiently low levels enables adoption within existing organizational constraints rather than requiring priority reordering.

Keywords: SME cybersecurity; cyber risk management; adoption barriers; non-adoption; conversational AI; voice interface; accessibility; organizational constraints

1 Introduction

Small and Medium Enterprises (SMEs) operate in an increasingly hostile cyber threat landscape. Recent data shows 43% of UK businesses experienced cyber security breaches or attacks in the last 12 months (DSIT, 2025), yet structured cyber risk management adoption remains critically low. 55% of SMEs lack defined cyber-risk strategies (Benz & Chatterjee, 2020), only 14% of microenterprises conduct ICT risk assessments (Sukumar et al., 2023), and exploratory survey data (n=16) collected for this research found a 25% adoption rate. This gap persists despite documented threats and mounting regulatory pressure from directives like the EU's NIS2 Directive (EU 2022/2555), which mandates cybersecurity risk management for organizations across critical sectors. With Ireland among 19 EU Member States currently finalizing transposition following European Commission enforcement action in May 2025 (European Commission, 2025), compliance urgency now extends beyond voluntary best practice.

Literature identifies five interconnected barriers preventing SME adoption. Time constraints, cost limitations, technical complexity, expertise gaps, and limited awareness (Heidt et al., 2019; Renaud & Ophoff, 2021; Kumar et al., 2024). Exploratory survey data collected for this research (n=16 SMEs) validates these barriers empirically whilst revealing that non-adoption persists despite awareness. Critically, survey findings reveal 89% recognize cybersecurity importance yet 56% cite competing priorities, indicating these technological barriers operate within organizational contexts of limited resources. This suggests barrier reduction must achieve thresholds low enough to make cybersecurity compatible with existing constraints rather than requiring priority reordering.

Existing solutions systematically fail SMEs. Comprehensive frameworks (NIST CSF, ISO 27001) impose resource burdens beyond SME capacity, simplified tools (CR-SAT, CISRAM) maintain cognitive burden through manual questionnaires, and emerging AI solutions remain theoretical or require unavailable technical infrastructure (Rawindaran et al., 2025; Heidt et al., 2019). Literature explicitly identifies this gap. Limited research explores how AI and ML can enhance cybersecurity accessibility for SME environments, with consistent calls for user-friendly interfaces addressing documented barriers (Rawindaran et al., 2025).

This research addresses three questions: (RQ1) Can voice-based conversational AI reduce technological adoption barriers sufficiently to enable adoption within existing organizational constraints? (RQ2) What accuracy trade-offs exist between manual and AI-assisted completion? (RQ3) Which AI model size optimizes cost-performance for practical deployment? This research employs mixed methods. Exploratory survey validates barriers and solution receptiveness, whilst technical evaluation demonstrates a voice-based conversational AI system guiding SMEs through CISRAM IG1 completion via natural language dialogue.

This research contributes empirically by validating literature-documented barriers in contemporary contexts (n=16 survey) and demonstrating substantial barrier reduction through voice-based conversational AI whilst maintaining operational accuracy. Methodologically, findings establish that system architecture determines performance alongside computational scale, with implications for LLM evaluation protocols in high-stakes domains.

The paper proceeds as follows: Section 2 synthesizes literature establishing SME barriers and solution failures; Section 3 outlines mixed-methods methodology; Sections 4-6 detail system design and implementation; Section 7 presents comparative evaluation results; Section 8 discusses implications, contributions, and limitations; Section 9 concludes with significance and future directions.

2 Literature Review

This review synthesizes 80+ sources from leading journals in information systems, cybersecurity, and human-computer interaction, using systematic search, backward citation tracking (Webster & Watson, 2002), and forward tracking via Web of Science. The literature encompasses diverse methodologies (i.e. surveys, structural equation modelling, qualitative interviews, and design science prototypes) providing triangulated evidence of SME barriers.

Theme 1: The SME Cybersecurity Paradox

SMEs face a critical paradox. Despite 39% identifying cyber-attacks in the last 12 months (Tsiodra et al., 2023), only 14% of microenterprises actively participate in ICT risk

assessments (Sukumar et al., 2023), and 55% lack any defined cyber-risk strategy (Benz & Chatterjee, 2020). This adoption gap stems from interconnected barriers spanning financial, technical, and human factors.

Resource constraints emerge as the primary obstacle. SMEs prioritize core business operations over IT security investments due to limited budgets (Heidt et al., 2019; Hoppe et al., 2021). Barrier quantification identifies skill gap and cost as strongest driving factors (Abdalla et al., 2025). Causality is empirically established. Lack of in-house technology skills has a direct and statistically significant relationship with readiness barriers ($\beta=0.458$, $p=0.000$) (Kumar et al., 2024).

Beyond resource constraints, cognitive overwhelm manifests as active avoidance behaviour. Empirical studies reveal 56% of SMEs report information overload, 55% believe compliance with all available advice is impossible, and critically, 47% actively avoid seeking cybersecurity information despite recognizing its importance (Renaud & Ophoff, 2021). Independent internet research on cybersecurity negatively impacts SME decision-making due to information overload (Salzberger, 2025). SME owners experience paralysis from low skill levels and lack of confidence about where to begin (Cartwright et al., 2023). This overwhelm translates to preparedness failures with 83% of SMEs lacking incident response plans entirely (Curtin et al., 2025), exposing critical vulnerability when attacks occur.

These literature-documented barriers receive empirical validation through primary survey data collected for this research. Among $n=16$ SME respondents across 7 industries, only 25% have adopted cyber risk management practices, with universal barrier recognition, 100% of respondents who ranked barriers cited all five factors (time, cost, complexity, expertise, awareness). Critically, 89% of non-adopters recognize cyber risk as 'very/extremely important' yet 56% cite 'other business priorities take precedence.' This extends beyond cognitive overwhelm to reveal a resource allocation challenge. Awareness of risk importance does not guarantee prioritization when competing demands (revenue generation, customer service, product development) offer more immediate returns. Time constraints thus reflect not just availability, but organizational priority decisions under resource scarcity.

Theme 2: Why Existing Solutions Fail SMEs

A variety of cybersecurity frameworks exist, yet these consistently fail to achieve widespread SME adoption due to fundamental misalignments with SME operational realities. A systematic review of cybersecurity maturity models identifies lack of end-to-end evaluation metrics, insufficient industry-specific focus, and limited coverage of emerging technologies like AI (Büyüközkan & Güler, 2025).

Existing solutions fail through three modes. First, comprehensive frameworks (NIST CSF, ISO 27001, Cyber Essentials) impose resource burdens beyond SME capacity through complexity, cost, time requirements, and extensive guidance requiring expert interpretation (Rawindaran et al., 2025; Heidt et al., 2019). Second, simplified SME-specific tools (CR-SAT, LCCI, CISRAM) maintain cognitive burden through manual questionnaires requiring 5+ hours completion for the uninitiated, creating the paralysis SMEs experience with information overload (Pawar & Palivela, 2022; Cartwright et al., 2023). Third, emerging AI-driven solutions remain theoretical without user-friendly implementation (ROHAN, CGF), require unavailable structured data (SME-SHIELD), or demand technical expertise for configuration (Multi-Fuzzy Inference Systems) (Asılı & Bahtiyar, 2024; Natsheh & Tabook, 2025).

The literature explicitly identifies this gap: limited research explores how AI and ML can mitigate cybersecurity threats in SME environments (Rawindaran et al., 2025). Researchers consistently call for developing user-friendly interfaces for ML-based solutions to enhance SME accessibility (Rawindaran et al., 2025).

Theme 3: Theoretical Foundations for Accessible AI Solutions

Artificial intelligence and conversational interfaces offer mechanisms to address the three primary barriers: expertise gaps, cognitive overwhelm, and resource constraints.

Machine learning demonstrates capability in threat detection and risk quantification, achieving substantial error reduction while maintaining expert judgment correlation (Ding et al., 2024; Ali et al., 2025; Natsheh & Tabook, 2025), suggesting AI can substitute for absent security specialists.

Conversational interfaces reduce cognitive overwhelm by eliminating complex navigation through natural language interaction (Ding et al., 2024; Alharbi et al., 2025). Cognitive Load Theory explains this as extraneous load from interface design reduces problem-solving capacity; conversational interfaces eliminate this burden (Surbakti et al., 2024).

Technology Acceptance Model demonstrates Perceived Ease of Use drives adoption (Davis, 1989), with voice interaction eliminating interface learning requirements. Diffusion of Innovation theory establishes lower complexity accelerates adoption (Rogers, 2003), supported by evidence that voice-integrated systems achieve 88.4% accessibility (Sai et al., 2025).

However, AI-driven approaches face limitations. Large Language Models exhibit hallucinations undermining reliability (Bhusal et al., 2024), black-box opacity complicates trust (Ding et al., 2024; Ali et al., 2024), and inherent biases affect accuracy (Afjal, 2024). User acceptance barriers exist, with voice interfaces assuming linguistic fluency and technological access not universally present (Malodia et al., 2022). These limitations necessitate rigorous evaluation against expert benchmarks.

Theme 4: The Research Gap and Novel Contribution

Documented SME barriers and systematic solution failures reveal a critical gap. No validated, framework-aligned, voice-based conversational AI solution exists for SME cyber risk assessment. A gap exists in systems beyond standard voice assistant style to evaluate novel and proactive conversational types (Cowan et al., 2022). Furthermore, previous studies predominantly focus on initial readiness issues, leaving a scarcity of inquiries examining actual implementation outcomes (Abdirad & Krishnan, 2020).

This research addresses these gaps by developing and evaluating a voice-based conversational AI tool that guides SMEs through CISRAM IG1 completion via natural language conversation. The system combines three elements currently separated in literature. Conversational voice modality for natural interaction, framework-aligned structure ensuring systematic coverage, and SME-focused design requiring minimal infrastructure. The evaluation compares voice-based completion against manual baseline across time, cost, and accuracy metrics, with secondary analysis of model sizing (20B vs 120B) to inform deployment decisions. This responds to calls for SME cybersecurity policy emphasizing systematic risk assessment over general awareness campaigns alone (Sukumar et al., 2023).

3 Research Methodology

This research adopts a pragmatic paradigm employing mixed methods to address the research gap, the absence of validated, framework-aligned, voice-based conversational AI solutions for SME cyber risk assessment. The methodology comprises two components. An exploratory survey research validates literature-documented barriers and establishes SME solution preferences, while technical evaluation demonstrates solution feasibility through comparative case study. The study is primarily quantitative, measuring barrier reduction across three dimensions identified in literature, time, cost, and complexity (Heidt et al., 2019; Renaud & Ophoff, 2021; Kumar et al., 2024) and assessing accuracy trade-offs between AI models.

Exploratory Survey

To empirically validate the barrier framework established in literature and assess SME receptiveness to simplified solutions, an exploratory survey was conducted between October and December 2025 targeting SME owners, IT Managers and decision-makers. The survey instrument comprised 20-23 questions depending on adoption status addressing current cyber risk management adoption status, perceived barriers to adoption (with ranked importance), budget allocation for cybersecurity, desired features in assessment tools, and willingness to adopt simplified alternatives.

Distribution occurred through multiple channels including LinkedIn professional networks and channels, WebSummit conference attendee groups, Facebook business communities, and relevant subreddits targeting SME decision-makers. The survey yielded n=16 responses from SMEs (excluding enterprises with 250+ employees) across 7 industries (Technology/IT Services, Manufacturing, Finance, Professional Services, Healthcare, Education, NGO) and 4 geographic regions (UK, US, EU, Australia/New Zealand). Ethics approval was obtained from the supervising academic prior to data collection.

Analysis employed descriptive statistics to establish adoption rates, barrier recognition frequencies, and solution preference distributions. The exploratory scope and small sample size limit statistical generalizability. However, findings validate barrier existence and solution receptiveness within this sample but cannot quantify population-level adoption likelihood without larger confirmatory studies. Survey responses informed technical system design decisions evaluated in the subsequent technical evaluation phase.

Research Questions

This research addresses three questions:

- RQ1: Can voice-based conversational AI reduce technological barriers (time, cost, complexity, expertise) sufficiently to enable SME adoption of cyber risk management practices within existing organizational constraints?
- RQ2: What is the accuracy trade-off between manual expert completion and AI-assisted completion of a CISRAM IG1 assessment?
- RQ3: Which AI model size (20B vs 120B parameters) provides the optimal cost-performance trade-off for practical SME deployment?

Technical Evaluation

The technical evaluation employed single-case methodology focusing on a 4-person EdTech micro-enterprise. The ground truth baseline consists of 56 safeguard maturity scores (1-3 scale) established through documented 5-6 hour manual CISRAM IG1 assessment. This

follows established practice in AI evaluation where automated systems are validated against expert judgment (Natsheh & Tabook, 2025).

Four evaluation approaches enabled comprehensive comparison. Manual Baseline (researcher's documented completion), AI-Assisted 20B (cost-optimized model), AI-Assisted 120B (accuracy-optimized model), and ChatGPT5.1 Baseline (unguided processing). Voice interview data capturing 51 CISRAM question responses served as consistent input across all AI approaches. Extended validation employed five runs per model (n=5) to account for LLM inference variance (Bhusal et al., 2024).

Metrics collected include processing time and cost (barrier reduction), Mean Absolute Deviation (MAD) comparing AI-generated scores against manual, and perfect match rate (safeguards scored identically). Detailed data collection procedures are described in Section 6 (Implementation).

4 Design Specification

This system transforms traditional manual CISRAM IG1 assessments into guided voice conversations, targeting 15-25 minute completion. User-selectable model architecture (20B vs 120B) directly enables RQ3 comparative evaluation of cost-performance trade-offs. The system comprises voice-driven client, LLM semantic analysis, workflow orchestrator, and report generation.

Architecture Design

The system implements a three-stage pipeline architecture (Figure 1):

Stage 1: Voice Capture and Transcription

Browser MediaRecorder API captures audio. Three-tier recognition ensures cross-browser compatibility. Groq Whisper Large v3 Turbo provides primary transcription, Browser Web Speech API serves as Chrome/Edge fallback, and Xenova Transformers enables offline Firefox/Safari functionality. Dead air detection triggers semantic analysis and advances the conversation after 10 seconds of silence.

Stage 2: Semantic Analysis and Question Extraction

Accumulated transcript is processed by user-selectable LLM for comparative evaluation (GPT-OSS-120B vs GPT-OSS-20B). Structured prompts extract which CISRAM questions were answered from conversational speech. Standardized JSON output enables systematic accuracy evaluation.

Four anti-hallucination mechanisms address specific failure modes documented in literature (Bhusal et al., 2024). Volume capping limits extraction to 10 answers per cycle, preventing over-helpful behaviour where models fabricate responses to maximize coverage. Confidence filtering ($\geq 3/5$ threshold) discards uncertain extractions, addressing the tendency to infer answers from tangentially related conversation. Duplicate detection prevents contradictory overwrites across analysis cycles, maintaining answer stability when users revisit topics. JSON repair handles truncated responses when token limits terminate generation mid-output, completing bracket structures through schema validation.

Stage 3: Report Generation

At the end of the conversation, the user reviews extracted answers in modal interface, enabling corrections before submission. HTTP POST to n8n webhook initiates two-stage workflow. Stage 1 employs Llama 3.3 70B to score maturity levels for 56 safeguards while

Stage 2 generates risk treatment plans. Pre-formatted Excel template is populated with scores and recommendations, then delivered to the user via SMTP email.

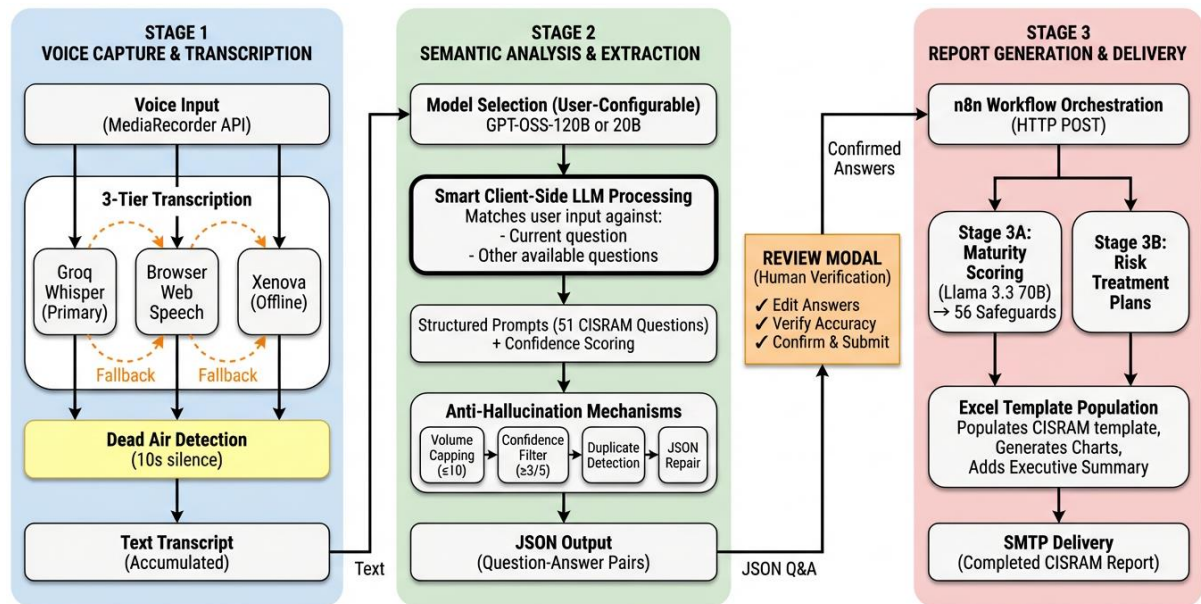


Figure 1: System Architecture - Three-Stage Voice-Based Assessment Pipeline

Technology Selection Rationale

Groq Whisper API Turbo chosen for superior accuracy on cybersecurity terminology and real-time latency. Multi-tier recognition addresses browser compatibility. User-selectable LLM models enable comparative evaluation for RQ3 cost-performance optimization. n8n workflow orchestration enables visual design for non-programmers.

Key Design Decisions

Dead air trigger (10 seconds) addresses conversational flow requirements. After detecting 10 seconds of silence, the system triggers semantic analysis to progress conversation. This duration allows natural thinking pauses without premature interruption. A 5-second countdown timer alerts users before automatic progression. Testing revealed shorter intervals (5-7 seconds) interrupted natural speech and caused unnecessary panic, while longer intervals (15+ seconds) created uncertainty about system responsiveness.

Confidence-based filtering ($\geq 3/5$ threshold) emerged from hallucination testing where lower thresholds ($\geq 2/5$) produced 40% false positive extractions. The midpoint threshold forces conservative extraction. The LLM must demonstrate moderate certainty, rejecting speculative inferences whilst maintaining recall for genuinely answered questions.

Two-stage workflow addresses token limit constraints and improves overall accuracy through separation into maturity scoring and treatment planning, improving reliability and enabling independent testing. Review modal verification ensures accuracy through human oversight, making AI decisions auditable before submission.

Design Alternatives Considered

Form-based UI was rejected due to higher completion times and high cognitive burden, consistent with documented SME avoidance of complex materials (Renaud & Ophoff, 2021; Cartwright et al., 2023). Immediate per-utterance triggering was rejected in favor of batched

analysis after conversational pauses (10s dead air), preserving context while managing API costs.

5 Implementation & Solution Development

This architecture was implemented over four months (September-December 2025) using an iterative development approach. The development environment comprised Visual Studio Code as primary IDE, Docker for n8n deployment, and React with Vite for quick frontend iteration. CISRAM IG1 was chosen for its SME-appropriate scope (56 safeguards, 51 assessment questions), representing foundational cybersecurity hygiene suitable for resource-constrained organizations.

Critical Implementation Issues

Issue 1: Keyword Matching vs Semantic Analysis

Initial implementation attempted to match user voice responses to CISRAM questions using keyword matching algorithms. Testing revealed this approach fundamentally inadequate as generic keywords failed to capture answer intent appropriately, missing contextual nuances in conversational speech. The impact was low question coverage and frequent misclassification of responses.

The solution migrated to LLM-based semantic analysis. The transcript is processed by GPT-OSS-120B or GPT-OSS-20B with structured prompts explicitly listing all 51 questions. The LLM continuously analyses conversational context in real-time, intelligently identifying which questions were implicitly or explicitly answered throughout natural dialogue. The system captures key insights as users speak naturally about their business, extracting relevant response segments with confidence scoring rather than requiring users to address questions sequentially. This real-time semantic understanding enables fluid conversation where users can volunteer information in any order, and the system automatically maps insights to appropriate framework controls.

Issue 2: Formula Preservation in Excel Templates

The n8n workflow required downloading the CISRAM IG1 Excel template, populating maturity scores and risk calculations, and returning a completed assessment. Initial implementation used SheetJS library for Excel manipulation. Testing revealed all formatting was stripped and embedded formulas broke during this process. Risk calculations depending on these formulas became unreliable, rendering templates unusable for compliance purposes.

The solution to resolve replicated all Excel formulas as JavaScript code within n8n workflow nodes. Maturity scores, expectancy calculations, and risk prioritization logic were reimplemented programmatically, maintaining calculation accuracy within tooling constraints. Formatting loss was accepted as tolerable for proof-of-concept demonstration, though production deployment would require commercial Excel manipulation libraries or alternative report formats.

Issue 3: Local LLM Performance Inadequacy

Initial backend implementation used Ollama running locally with llama3.2:3b for the n8n workflow stage that generates maturity scores from extracted question-answer pairs. Performance testing revealed processing times consistently exceeding 10 minutes per report generation, making iterative development and testing extremely time-consuming. Additionally, the 3B model produced inconsistent accuracy across test runs, with maturity scores varying significantly for identical inputs.

Migration to Groq's cloud API using llama-3.3-70b-versatile reduced processing time from 10+ minutes to 2-4 seconds whilst substantially improving accuracy through the 70-billion parameter model's enhanced context understanding. The cost per assessment (€0.02-€0.04 API charges) proved negligible compared to development velocity gains and accuracy improvements, addressing cost constraints identified as a primary adoption barrier (Heidt et al., 2019; Abdalla et al., 2025).

Issue 4: Single-Stage Token Limit Constraints

Initial n8n workflow design attempted to generate complete CISRAM reports (56 safeguard maturity scores plus detailed risk treatment recommendations with implementation guidance) in a single LLM API call. Testing revealed this approach consistently exceeded Llama 3.3 70B's free-tier 8192 token context window. Outputs were either incomplete (truncated mid-generation) or malformed (missing closing brackets, invalid JSON structure).

The solution restructured the workflow into two independent stages with separate LLM calls. Stage 1 processes the 51 extracted question-answer pairs and generates maturity scores (1-3 scale) for all 56 safeguards. Stage 2 generates risk treatment plans. This separation improved reliability dramatically, enabled independent unit testing of each stage, and maintained all outputs within token constraints.

Issue 5: LLM Hallucination and Inference Variance

Semantic analysis testing revealed the LLM generated false answer extractions, confirming hallucination concerns documented in literature (Bhusal et al., 2024), particularly when attempting to maximize question coverage to appear helpful. The model would occasionally invent plausible-sounding responses for questions users had not addressed, conflating information from different conversational contexts.

Four mechanisms address reliability concerns. Three mechanisms specifically target hallucination by forcing conservative extraction. Volume capping limits extraction to 10 answers per cycle, confidence filtering ($\geq 3/5$ threshold) discards uncertain extractions, and duplicate detection maintains answer stability across analysis cycles. JSON repair addresses a separate technical failure mode, correcting truncated LLM responses through bracket completion and schema validation when generation terminates mid-output. The technical lesson here shows single-run LLM evaluation is fundamentally unreliable due to inference variance; multi-run aggregation ($n \geq 5$) provides stable performance estimates, and architectural constraints (volume caps, confidence thresholds) are essential for production reliability.

Issue 6: Transcription Accuracy with Technical Terminology

Initial voice interface implementation used browser SpeechRecognition API for real-time audio-to-text transcription. Testing with actual cybersecurity terminology revealed insufficient accuracy, with technical terms frequently misrecognized or transcribed phonetically incorrectly. These transcription errors propagated through the semantic analysis pipeline with incorrect terminology in transcripts led to question-answer mapping failures, as the LLM could not match garbled technical terms to CISRAM control definitions.

Solution involved upgrading to Groq Developer tier, gaining access to Whisper Large v3 Turbo for speech-to-text processing. This improved transcription quality for specialized cybersecurity terminology, addressing the misrecognition issues observed with browser-native speech recognition. The upgrade simultaneously addressed throughput constraints by

increasing rate limits to support continuous processing demands. Cost-accuracy trade-off analysis showed €0.02-€0.04 per audit represented negligible expense for substantial accuracy improvement and system reliability gains. The technical lesson here is that domain-specific applications with specialized terminology require purpose-built transcription models rather than general-purpose browser APIs; accuracy improvements in early pipeline stages compound throughout system, justifying API costs for quality-critical components.

System Functionality Validation

The final implementation was validated through functional testing on Windows desktop environment. The primary transcription tier (Groq Whisper Large v3 Turbo) functioned consistently across Chrome, Edge, and Firefox browsers. The system successfully completed 10 full end-to-end assessment cycles (5 runs $\times$ 2 models) without requiring manual intervention or error recovery.

Validation testing confirmed anti-hallucination mechanisms reduced false extractions. Volume capping prevented over-extraction, confidence filtering rejected uncertain responses, duplicate detection maintained answer stability, and JSON repair handled three truncated responses successfully.

The implementation produced evaluable datasets suitable for comparative analysis. Time measurements were captured via timestamped recordings from the moment the record button was pressed to email delivery receipt (end-to-end process time). API costs were tracked through Groq billing records for all transcription and model inference calls. Accuracy data was generated through $n=5$ validation runs per model, producing maturity scores for 56 safeguards in each run. These scores were compared against the manual baseline to enable accuracy evaluation. These datasets enable the detailed comparative analysis presented in Section 7.

Outputs Produced

Time data shows processing durations of 13.4 minutes (20B model) and 12.95 minutes (120B model) captured via timestamped recordings, compared against manual baseline of 4-6 hours for new adopters. Cost data shows API expenses of €0.03 (20B) and €0.04 (120B) per assessment tracked through Groq billing records, compared against manual opportunity costs of €200-250 (assuming a €50.00 per hour rate). Accuracy data shows maturity scores (1-3 scale) for 56 safeguards across $n=5$ runs per model, generating comparison datasets for Mean Absolute Deviation calculation and perfect match rate analysis.

Implementation Limitations

The current implementation has technical and methodological constraints requiring attention for production deployment.

Third-party API dependency on Groq services for transcription and semantic analysis raises data sovereignty concerns for regulated industries. Production deployment requires enterprise data processing agreements or migration to self-hosted model deployment with sufficient computational resources. The current cloud-based architecture prevents real-world SME testing with sensitive business data, as participants would need to disclose confidential operational information through third-party services. This limitation reduced the study to single-case technical evaluation rather than multi-organization field testing, constraining generalizability of findings.

Formula preservation using JavaScript replication reduces maintainability compared to native Excel formulas; commercial Excel manipulation libraries would restore formula fidelity. Browser compatibility testing was limited to desktop environments (Windows) on three major browsers (Chrome, Edge, Firefox). Comprehensive mobile validation and cross-platform testing (Mac, Linux) remain incomplete. The lack of production deployment infrastructure prevented longitudinal evaluation of actual adoption patterns and sustained usage behaviour, limiting findings to initial feasibility demonstration rather than real-world adoption validation.

6 Evaluation

This evaluation tests the voice-based system against three research questions using quantitative metrics (time, cost, accuracy) and survey validation as outlined in Section 3.

Survey Results: Empirical Problem Validation

Prior to technical evaluation, survey data from n=16 SME owners established baseline adoption rates and barrier perception, validating the problem statement and solution design rationale.

Metric	Finding
Sample size	n=16 SMEs (excluding 250+ employees)
Industries	7 sectors (Tech/IT, Manufacturing, Finance, Professional Services, Healthcare, Education, NGO)
Geographic distribution	UK, US, EU, Australia/New Zealand
Adoption rate	25% (4/16 SMEs)
Started but incomplete	19% (3/16 SMEs)

Table 1: Survey Demographics

Adoption Rate: Confirming the Problem

Only 25% of surveyed SMEs (4/16) have successfully adopted cyber risk management practices. An additional 19% (3/16) initiated adoption but failed to complete implementation, revealing that barriers persist even after initial commitment. Combined with 56% (9/16) who never attempted adoption, this confirms that three-quarters of SMEs lack functional cyber risk management frameworks despite documented threats.

Barrier Validation: Universal Acknowledgment

Among respondents who ranked barriers (n=9), 100% cited all five factors, indicating universal barrier recognition:

Barrier	Specific Evidence
Time constraints	"Don't have the time; Other business priorities take precedence"
Complexity	Requested "simpler, less technical template"
Technical expertise	Requested "professional guidance/consulting"
Cost/budget	89% have no dedicated cybersecurity budget
Awareness	78% unaware of free resources (CISRAM, NCSC)

Table 2: Barrier Recognition Rates

These findings triangulate with literature-documented barriers (Heidt et al., 2019; Renaud & Ophoff, 2021; Kumar et al., 2024), providing empirical validation in 2025.

Solution Validation: Addressing Stated Needs

Non-adopters' prioritized solution features directly align with the voice-based system architecture:

SME Need	Frequency	System Response
Faster process (<1hr)	78% (7/9)	13-minute completion (RQ1)
Simpler template	56% (5/9)	No template - conversational interface
Step-by-step guidance	56% (5/9)	Voice-guided question flow
Technical simplification	100% (9/9)	Plain-language prompts

Table 3: Solution Needs vs System Design

Critically, 78% of non-adopters indicated they would be "Likely" or "Very Likely" to adopt a simplified tool requiring under 1 hour, suggesting the adoption gap is technologically addressable rather than persistent behavioural resistance.

Critical Nuance: Competing Priorities vs Awareness

Survey findings complicate the "awareness" barrier. While 89% of non-adopters rated cyber risk as 'very/extremely important,' 56% cited 'other business priorities take precedence' as the reason for non-adoption. This distinction reveals technological barriers must be reduced to thresholds compatible with existing organizational constraints. When cybersecurity requires 5-6 hours and €200-250, it competes directly with urgent demands and loses despite recognized importance. Barrier reduction to 13 minutes and €1 makes cybersecurity light enough to coexist with competing priorities rather than compete.

Technical Results: RQ1 - Barrier Reduction

Time Barrier

Approach	Duration	Key Context
Manual (Baseline)	4-6 hours	Initial unfamiliar CISRAM completion
20B POC	13.4 minutes	With transcript reading
120B POC	12.95 minutes	With transcript reading

Table 4: Time Comparison Across Assessment Approaches

Table 4 demonstrates that the voice-based system reduced assessment time from 4-6 hours to approximately 13 minutes for both models, representing a 24-fold improvement over the baseline duration. For experienced practitioners conducting subsequent assessments, the improvement remains substantial at 12-18 times faster than manual completion.

Cost Barrier

Approach	Direct Cost	Opportunity Cost	Total Cost
Manual (SME staff)	€0 (internal)	€200-250	€200-250
20B POC	€0.03 API	€0.86	~€1
120B POC	€0.04 API	€0.86	~€1

Table 5: Cost Comparison per Assessment

Table 5 shows the system reduces per-assessment cost from €200-250 to approximately €1, representing over 200-fold cost reduction. The cost premium between models proves negligible. The 33% relative difference (€0.03 vs €0.04) represents €0.01 in absolute terms per assessment.

Knowledge Barrier

The manual CISRAM process requires users to navigate technical documentation, understand maturity frameworks, and cross-reference complex assessment criteria. The system eliminates this requirement through guided conversation. Users answer plain-language questions about their business rather than interpreting technical control definitions, with the system mapping responses to safeguard assessments internally. The voice-based system substantially reduces all three barriers (time, cost, knowledge).

Technical Results: RQ2 - Accuracy Trade-off

Extended Validation: Addressing Inference Variance

Initial single-run testing suggested 20B model superiority with lower Mean Absolute Deviation (MAD 0.2500) compared to 120B (MAD 0.2857). However, this finding proved misleading when subjected to extended validation.

Metric	20B POC	120B POC	ChatGPT5.1
MAD Score	0.2321	0.1250	0.3929
MAD %	11.61%	6.25%	19.64%
Perfect Matches	43/56 (76.8%)	49/56 (87.5%)	35/56 (62.5%)
Off by 1 Point	13/56 (23.2%)	7/56 (12.5%)	20/56 (35.7%)
Off by 2+ Points	0/56 (0%)	0/56 (0%)	1/56 (1.8%)

Table 6: Accuracy Metrics Comparison Against Manual Baseline (n=5 runs per model)

Extended testing with five runs per model completely reversed the initial finding. Comparing AI-generated scores against the manual baseline, Table 6 shows 120B achieved approximately 6% Mean Absolute Deviation versus 20B's 12%, representing 86% improvement in accuracy. The 120B model produced 49 perfect matches compared to 20B's 43, with only 7 safeguards differing by one point and zero exceeding one-point deviation. Multi-run aggregation (n=5) provides stable performance estimates, revealing 120B's consistent superiority across iterations.

Error Pattern Analysis

The 120B model's errors demonstrate conservative bias. When deviating from manual baseline, the system underestimated maturity rather than overestimated. Of the 7 one-point deviations, 5 rated safeguards lower than manual assessment. The 20B model exhibited systematic bias on infrastructure controls, particularly network firewall assessment. Manual assessment scored this as "1" (ad-hoc implementation); 20B persistently scored "2" across all five runs, representing false positive. The 120B model correctly scored "1" in all runs.

Architecture Validation: ChatGPT5.1 Comparison

The unguided ChatGPT5.1 baseline achieved MAD of approximately 20% error with only 35 perfect matches, representing over 3-fold accuracy disadvantage compared to 120B. ChatGPT5.1 received identical source documents (CISRAM framework, safeguard definitions, IG1 template) and the same business transcript, yet achieved substantially lower accuracy. This comparison demonstrates that system architecture (i.e. guided extraction,

structured prompting, review modal verification) contributes more to accuracy than model selection alone. The 120B model achieves 87% perfect match rate with acceptable 6% error for operational use.

Technical Results: RQ3 - Model Sizing

Dimension	20B	120B	Trade-off
Cost per Assessment	€0.03	€0.04	+33% premium
Processing Speed	13.4 min	12.95 min	Equivalent
Accuracy (MAD)	0.2321	0.1250	120B 1.86x better
System Stability	3 errors	2 errors	120B superior

Table 7: Model Comparison – 20B vs 120B

Table 7 demonstrates that 120B costs €0.04 per assessment versus 20B's €0.03, representing 33% relative premium but only €0.01 absolute difference. Against this negligible cost increase, 120B delivers approximately 86% better accuracy, 11 percentage point improvement in perfect match rate, and superior system stability. The 120B model is optimal despite 33% cost premium due to negligible absolute cost difference (€0.01) and substantial accuracy gains.

7 Discussion

This research addressed three questions investigating whether voice-based conversational AI could reduce SME cyber risk management adoption barriers, what accuracy trade-offs exist, and which model size optimizes deployment. Results demonstrated substantial barrier reduction across time (24-fold), cost (200-fold), and knowledge requirements, with the 120B model achieving 87% accuracy through architectural design rather than computational scale alone, directly addressing the problem statement's identified adoption gap.

Validation of Literature Gaps

The evaluation validates literature's explicit call for exploration of how artificial intelligence and machine learning can enhance and mitigate the cybersecurity landscape in SME environments (Rawindaran et al., 2025), providing empirical evidence that voice-based conversational AI addresses documented barriers. Survey data (n=16) empirically validates these barriers in contemporary contexts. The 25% adoption rate confirms documented non-adoption patterns, while universal barrier recognition (100% of rankers cite all five barriers) triangulates with literature findings (Heidt et al., 2019; Renaud & Ophoff, 2021; Kumar et al., 2024).

Critically, survey findings reveal a distinction between awareness and action. 89% of non-adopters recognize cyber risk management as "very/extremely important," yet 56% cite "other business priorities take precedence." This validates the organizational adoption gap. SMEs understand importance but deprioritize cybersecurity when competing demands arise. This extends Technology Acceptance Model by revealing its organizational limitation. Improved ease of use and usefulness affect individual acceptance but not organizational adoption when competing priorities intervene. However, survey findings suggest competing priorities are not absolute resistance but conditional on resource requirements. The 78% expressing willingness to adopt tools requiring under one hour validates that technological barrier reduction to sufficiently low thresholds enables adoption within existing constraints, making cybersecurity 'light enough to fit', rather than requiring priority reordering. The

distinction between technological feasibility (demonstrated through barrier reduction) and adoption likelihood (requiring organizational readiness) represents a critical gap requiring behavioural research beyond technological intervention.

Extension of Existing Frameworks

This research extends literature through three contributions. First, the architecture-over-model principle challenges AI industry narratives emphasizing computational scale. The 3-fold accuracy advantage over ChatGPT5.1 despite frontier model capabilities demonstrates that system design (i.e. guided extraction, structured prompting, review verification) determines performance more than raw model power. This suggests interface engineering and workflow architecture yield greater returns than exclusively pursuing larger parameter models.

Second, the single-run to multi-run reversal establishes that LLM evaluation requires $n \geq 5$ validation protocols. Single-test metrics capture noise rather than capability.

Third, systematic bias characterization reveals model sizing affects error patterns, not just accuracy magnitude. The 20B model's persistent false positive on infrastructure controls (scoring "2" across all five runs when ground truth was "1") indicates systematic directional bias. In contrast, the 120B model's errors exhibited random variation with conservative underestimation. This distinction matters for deployment. Random errors are manageable through human review, but systematic biases require architectural mitigation or model replacement. Practitioners should consider error character alongside aggregate accuracy when selecting models for domain-specific tasks.

Stakeholder Implications

For SME practitioners, this demonstrates cognitive barriers are technologically addressable, transforming documented obstacles into managed constraints, though organizational commitment to act on findings remains essential. For academic researchers, the work validates voice modality for cyber risk assessment and highlights organizational adoption factors as critical research gaps, with methodological contributions offering guidance for LLM evaluation in high-stakes domains. For security tool vendors, the architecture principle suggests competitive advantage derives from interface design and workflow engineering as much as model selection, contradicting assumptions that accuracy requires primarily larger models.

Limitations & Generalizability

This proof-of-concept demonstrates technological feasibility within explicit constraints. The survey's exploratory scope ($n=16$) validates barrier existence but cannot establish population-level generalizability without larger confirmatory studies. Single-case technical evaluation limits extrapolation to broader SME populations, different industries, or alternative compliance frameworks. Critically, this research demonstrates technological barriers are addressable but does not claim the solution drives adoption. Organizational factors (management prioritization, perceived threat salience, competing demands) remain unaddressed. The system enables assessment completion by removing certain barriers. However, it does not guarantee organizations act on findings.

Future Research Directions

Immediate priorities include multi-organization testing ($n \geq 25$ diverse SMEs across industries and regions) to establish generalizability and comparison against professional consultants to quantify accuracy gaps for compliance reporting. User acceptance testing with naive SME

owners would validate whether barrier reduction translates to actual adoption among target populations. Critical research questions emerge from the organizational adoption gap. Investigating why SMEs fail to adopt despite removed technological barriers requires behavioural methodologies beyond this study's scope. Longitudinal studies measuring whether AI-generated assessments drive actual security improvements would validate practical value beyond assessment completion.

8 Conclusion

This research investigated whether voice-based conversational AI could address the persistent SME cyber risk management adoption gap, specifically examining barrier reduction, accuracy trade-offs, and optimal model sizing for practical deployment.

The findings establish technological feasibility for addressing the documented adoption crisis where only 25% of SMEs implement frameworks despite 43% experiencing attacks annually (DSIT, 2025). Voice-based assessment achieves 24-fold time reduction (5-6 hours compared to 13 minutes) and 200-fold cost reduction (€200 compared to €1) whilst maintaining 87% accuracy, demonstrating that documented barriers (time, cost, complexity, expertise) are technologically addressable. However, survey evidence reveals the critical limitation and opportunity. 89% recognize cybersecurity importance yet cite competing priorities, indicating technological barriers must reach thresholds enabling coexistence with competing demands. The 78% expressing willingness to adopt tools requiring under one hour validates this threshold exists. Barrier reduction to 13 minutes and €1 makes cybersecurity light enough to fit within organizational constraints rather than compete directly with urgent operational demands.

This research contributes three insights beyond feasibility demonstration. First, system architecture determines performance more than computational scale. The structured extraction approach achieved 3-fold accuracy advantage over unstructured ChatGPT5.1 baseline (19.6% vs 6.25% MAD), demonstrating that system architecture (guided extraction, structured prompting, review verification) determines performance as much as model selection. Second, LLM evaluation requires $n \geq 5$ validation protocols, as single-run testing captured measurement noise rather than true capability. Third, model sizing affects error patterns. Smaller models exhibit systematic directional bias whilst larger models produce random errors, a distinction critical for deployment decisions in high-stakes domains.

These findings matter beyond SME cybersecurity. The architecture-over-model principle applies to any domain deploying LLMs for structured extraction tasks. The organizational adoption gap, where improved ease-of-use proves insufficient when competing priorities dominate, extends Technology Acceptance Model by revealing its limitation in resource-constrained environments. The convergence of regulatory pressure (NIS2 mandates) and technological accessibility may finally enable compliance where voluntary adoption failed, representing a natural experiment in policy-technology interaction worth monitoring.

Future research must shift from technological feasibility to organizational adoption mechanisms, investigating why simplified tools remain unused despite reduced barriers. The SME cybersecurity problem is solvable technologically. However, solving it practically requires understanding that organizations are not rational actors optimizing security. They are resource-constrained entities constantly triaging competing urgent demands. Yet the convergence demonstrated here, where 78% express adoption willingness when barriers are

removed, coupled with regulatory mandates like NIS2, suggests the persistent gap between SME vulnerability and protective action may finally be closing.

9 References

- ®) and HALOCK Security Labs, CIS Risk Assessment Method (RAM) now simplifies your path toward reasonable security with evidence-based guidance. Reasonable Risk Mission Risk Acceptable Risk Objectives Obligations CIS RAM V2. (n.d.).
- Abdalla, A.A., Tiwari, C.K., Bhat, Mohd.A., Pal, A. and Abdulmuhsin, A.A. (2025). Examining the impediments to AI chatbots adoption among professional practitioners: An ISM – SEM approach. *Acta Psychologica*, [online] 260, p.105488. doi:<https://doi.org/10.1016/j.actpsy.2025.105488>.
- Abdirad, M. and Krishnan, K. (2020). Industry 4.0 in Logistics and Supply Chain Management: A Systematic Literature Review. *Engineering Management Journal*, 33(3), pp.1–15. doi:<https://doi.org/10.1080/10429247.2020.1783935>.
- Afjal, M. (2024). Evolving trends, limitations, and ethical considerations in AI-driven conversational interfaces: assessing ChatGPT’s impact on healthcare, financial services, and educational sectors. *Technology Analysis & Strategic Management*, pp.1–20. doi:<https://doi.org/10.1080/09537325.2024.2420617>.
- Alharbi, N., Din, F.U., Paul, D. and Sadgrove, E. (2025). Driving AI Chatbot Adoption: A Systematic Review of Factors, Barriers, and Future Research Directions. *Journal of Open Innovation: Technology, Market, and Complexity*, [online] p.100590. doi:<https://doi.org/10.1016/j.joitmc.2025.100590>.
- Ali, S.M., Razzaque, A., Abbass, H., Yousaf, M. and Ali, S.S. (2025). A novel AI-Based Integrated Cybersecurity Risk Assessment Framework and resilience of National critical infrastructure. *IEEE Access*, pp.1–1. doi:<https://doi.org/10.1109/access.2024.3524884>.
- Asılı, H. and Bahtiyar, Ş. (2025). SME-SHIELD: A Scalable AI Architecture for Third-Party Cyber Risk Assessment in Small and Medium Enterprises (SMEs). pp.1374–1379. doi:<https://doi.org/10.1109/ubmk67458.2025.11206972>.
- Benz, M. and Chatterjee, D. (2020). Calculated risk? A cybersecurity evaluation tool for SMEs. *Business Horizons*, 63(4), pp.531–540. doi:<https://doi.org/10.1016/j.bushor.2020.03.010>.
- Bhusal, D., Alam, M.T., Nguyen, L., Mahara, A., Lightcap, Z., Frazier, R., Fieblinger, R., Torales, G.L. and Rastogi, N. (2024). SECURE: Benchmarking Generative Large Language Models for Cybersecurity Advisory. *arXiv (Cornell University)*. doi:<https://doi.org/10.48550/arxiv.2405.20441>.
- Büyüközkan, G. and Güler, M. (2025). Cybersecurity maturity model: Systematic literature review and a proposed model. *Technological Forecasting and Social Change*, 213, p.123996. doi:<https://doi.org/10.1016/j.techfore.2025.123996>.
- Cartwright, A., Cartwright, E. and Edun, E. (2023). Cascading information on best practice: Cyber security risk management in UK micro and small businesses and the role of IT companies. *Computers & Security*, 131, p.103288. doi:<https://doi.org/10.1016/j.cose.2023.103288>.

- Chandani, A., Agrawal, R., Ubarhande, P., Kumar, K. and Alhamzi, K.H.M. (2025). Exploring adoption barriers of ICT in SMEs using an ISM-MICMAC approach. *Humanities and Social Sciences Communications*, [online] 12(1). doi:<https://doi.org/10.1057/s41599-025-05295-9>.
- CISRAM CORE (n.d.). *CIS Risk Assessment Method (RAM) Version 2.1 CIS Risk Assessment Method (RAM) Background and Acknowledgments*.
- CISRAM IG1 (2022). *CIS Risk Assessment Method (RAM) Version 2.1 CIS Risk Assessment Method (RAM) Implementation Group 1 (IG1) Workbook Edition Background and Acknowledgments*.
- Cowan, B.R., Clark, L., Candello, H. and Tsai, J. (2023). Introduction to this special issue: guiding the conversation: new theory and design perspectives for conversational user interfaces. *Human-Computer Interaction*, pp.1–8. doi:<https://doi.org/10.1080/07370024.2022.2161905>.
- Curtin, M., Sheehan, B., Gruben, M., O’Carroll, G. and Murray, H. (2025). Enhancing Cybersecurity Awareness in Small and Medium Enterprises Through a User-Friendly Risk Assessment Tool. [online] pp.209–226. doi:<https://doi.org/10.1109/eurosp63326.2025.00021>.
- Davis, F.D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, [online] 13(3), pp.319–340. doi:<https://doi.org/10.2307/249008>.
- Ding, W., Abdel-Basset, M., Ali, A.M. and Moustafa, N. (2024). Large language models for cyber resilience: A comprehensive review, challenges, and future perspectives. *Applied Soft Computing*, [online] 170, p.112663. doi:<https://doi.org/10.1016/j.asoc.2024.112663>.
- Directive (EU) 2022/2555 (2025). *EUR-Lex - 32022L2555 - EN - EUR-Lex*. [online] eur-lex.europa.eu. Available at: <https://eur-lex.europa.eu/eli/dir/2022/2555/oj>.
- DSIT (2025). *Cyber security breaches survey 2025*. [online] GOV.UK. Available at: <https://www.gov.uk/government/statistics/cyber-security-breaches-survey-2025/cyber-security-breaches-survey-2025>.
- European Commission (2025). *State-of-play of the transposition of the NIS Directive | Shaping Europe’s digital future*. [online] digital-strategy.ec.europa.eu. Available at: <https://digital-strategy.ec.europa.eu/en/policies/nis-transposition>.
- Heidt, M., Gerlach, J.P. and Buxmann, P. (2019). Investigating the Security Divide between SME and Large Companies: How SME Characteristics Influence Organizational IT Security Investments. *Information Systems Frontiers*, 21(6), pp.1285–1305. doi:<https://doi.org/10.1007/s10796-019-09959-1>.
- Hoppe, F., Gatzert, N. and Gruner, P. (2021). Cyber risk management in SMEs: insights from industry surveys. *The Journal of Risk Finance*, ahead-of-print(ahead-of-print). doi:<https://doi.org/10.1108/jrf-02-2020-0024>.
- Kumar, R., Dutta, G. and Rakesh Kumar Phanden (2024). Digitalization Adoption Barriers in the Context of Sustainability and Operational Excellence: Implications for SMEs. *Engineering Management Journal*, pp.1–17. doi:<https://doi.org/10.1080/10429247.2024.2372519>.
- Malodia, S., Kaur, P., Ractham, P., Sakashita, M. and Dhir, A. (2022). Why do people avoid

and postpone the use of voice assistants for transactional purposes? A perspective from decision avoidance theory. *Journal of Business Research*, 146, pp.605–618. doi:<https://doi.org/10.1016/j.jbusres.2022.03.045>.

- Natsheh, E. and Tabook, F.B. (2025). AI-enhanced Cybersecurity Risk Assessment with Multi-Fuzzy Inference. *Journal of ICT Research and Applications*, 19(1), pp.1–26. doi:<https://doi.org/10.5614/itbj.ict.res.appl.2025.19.1.1>.
- Pawar, S. and Palivela, Dr.H. (2022). LCCI: A framework for least cybersecurity controls to be implemented for small and medium enterprises (SMEs). *International Journal of Information Management Data Insights*, 2(1), p.100080. doi:<https://doi.org/10.1016/j.jjime.2022.100080>.
- Rawindaran, N., Jayal, A. and Prakash, E. (2025). Cybersecurity Framework: Addressing Resiliency in Welsh SMEs for Digital Transformation and Industry 5.0. *Journal of Cybersecurity and Privacy*, [online] 5(2), p.17. doi:<https://doi.org/10.3390/jcp5020017>.
- Rawindaran, N., Jayal, A., Prakash, E. and Hewage, C. (2023). Perspective of Small and Medium Enterprise (SME's) and their Relationship with Government in Overcoming Cybersecurity Challenges and Barriers in Wales. *International Journal of Information Management Data Insights*, [online] 3(2), pp.100191–100191. doi:<https://doi.org/10.1016/j.jjime.2023.100191>.
- Renaud, K. and Ophoff, J. (2021). A cyber situational awareness model to predict the implementation of cyber security controls and precautions by SMEs. *Organizational Cybersecurity Journal: Practice, Process and People*, 1(1), pp.24–46. doi:<https://doi.org/10.1108/ocj-03-2021-0004>.
- Rogers, E.M. (2003). *Diffusion of innovations*. 5th ed. New York: Simon and Schuster.
- Sai, C.C., Raguraman Purushothaman, Reddy, C., Peddi Reddy Gangothri, V. Dhanush and Chintamanipeta Bhavana (2025). Integrating Optical Characteristic Recognition with Conversational AI: A Multimodal Chatbot Featuring Speech and Poster Generation. [online] pp.1522–1529. doi:<https://doi.org/10.1109/icaiss61471.2025.11041931>.
- Salzberger, A. (2025). An empirical analysis of the behavioral influences and information sources affecting the cyber insurance decisions of German SMEs. *The Journal of Risk Finance*. doi:<https://doi.org/10.1108/jrf-05-2024-0151>.
- Sukumar, A., Mahdiraji, H.A. and Jafari-Sadeghi, V. (2023). Cyber risk assessment in small and medium-sized enterprises: A multilevel decision-making approach for small e-tailors. *Risk Analysis*, 43(10). doi:<https://doi.org/10.1111/risa.14092>.
- Surbakti, R., Umboh, S.E., Pong, M. and Dara, S. (2024). Cognitive Load Theory: Implications for Instructional Design in Digital Classrooms. *International Journal of Educational Narratives*, [online] 2(6), pp.483–493. doi:<https://doi.org/10.70177/ijen.v2i6.1659>.
- Tsiodra, M., Panda, S., Chronopoulos, M. and Panaousis, E. (2023). Cyber Risk Assessment and Optimisation: A Small Business Case Study. *IEEE Access*, [online] 11, pp.1–1. doi:<https://doi.org/10.1109/ACCESS.2023.3272670>.
- Watson, R.T. and Webster, J. (2020). Analysing the past to prepare for the future: Writing a literature review a roadmap for release 2.0. *Journal of Decision Systems*, 29(3), pp.129–147. doi:<https://doi.org/10.1080/12460125.2020.1798591>.