

Table of Contents

Title	Page
Executive Summary	00
Chapter 1 - Introduction	01
1.1 Rationale	01
1.2 Target outcomes	02
1.3 Ground rules	02
Chapter 2 - Pilot participants and use cases	03
2.1 Participant profile	03
2.2 Use cases	04
2.3 Patterns of LLM usage	05
Chapter 3 - Risk Assessment and Test Design	06
3.1 Risk Assessment	06
3.2 Metrics	07
3.3 Testing approach: Test datasets	07
3.4 Testing approach: Evaluators	08-09
Chapter 4 - Test Implementation	10
4.1 Test Environment	10
4.2 Test data and effort	10
4.3 Implementation challenges	10
Chapter 5 - Lessons learnt	11
5.1 Test what matters	11-12
5.2 Don't expect test data to be fit for purpose	13
Guest Blog: Learning from self-driving cars: Simulation Testing	14
Guest Blog: Synthetic Data for Adversarial Testing	15
5.3 Look under the hood	16
5.4 Use LLMs as judges, but with skill and caution	17
Guest Blog: LLM-as-a-judge: Pros and Cons	18
5.5 Keep your human SMEs close!	19
Guest Blog: LLMs can't read your mind	20
Chapter 6 - What's next?	21-22

Executive Summary

From Model Safety to Application Reliability

As Generative AI (“GenAI”) transitions from personal productivity tools and consumer-facing chatbots into real-world environments like hospitals, airports and banks, it faces a higher bar on quality and confidence.

- 01 Risk assessments depend heavily on the **context** of the use case – e.g., lower tolerance for error in a clinical application than a customer service chatbot.
- 02 Given the higher **complexity** involved in integrating foundation models with existing data sources, processes and systems, there are more potential points of failure.

However, much of the current work around AI testing focuses on the *safety of foundation models*, rather than the *reliability of end-to-end applications*. The Global AI Assurance Pilot was an attempt to address this gap: not through academic research, but by building upon real-life experiences of practitioners around the world.

Learning by doing

The pilot matched **17 deployers** of GenAI applications with **16 specialist AI testing firms**. These organisations were based in Singapore and 8 other geographies, providing a significant **international** lens. The primary objective was to surface and codify emerging norms in technical testing of GenAI applications. The 17 applications were aimed at a mix of internal (12) and external (5) **users**. There was a **human in the loop** (12) in most cases. 10 **industries** were represented, including banking, healthcare and technology. Large Language Models (LLMs) were **utilised in a variety of ways** in these applications: summarisation, retrieval augmented generation, data extraction, chatbots, classification, translation, agentic workflows and code generation.

The “what” and “how” of testing GenAI applications



Deciding what to test (or not!)

Deciding what to test (or not!) was a non-trivial exercise. The 3 risks that interested most deployers were accuracy and robustness, use case specific regulation and compliance requirements, and content safety.



Test Datasets

Off-the-shelf LLM benchmark **test datasets** were rarely used to conduct

Getting GenAI testing right: 4 practical recommendations

- 01 Test what matters

Your context will determine what risks you should (or should not!) care about. Spend time upfront to design effective tests for those.
- 02 Don't expect test data to be fit for purpose

No one has the "right" test dataset to hand. Human and AI effort is needed to generate realistic, adversarial and edge case test data.
- 03 Look under the hood

Testing just the outputs may not be enough. Interim touchpoints in the application pipeline can help with debugging and redundancy. With agentic AI applications, this becomes a necessity.
- 04 Use LLMs as judges, but with skill and caution

Human-only evaluations will not scale. LLMs-as-judges are necessary but require careful design and human calibration. Cheaper, faster and simpler alternatives exist, in some situations.

There was also an overwhelming reinforcement of the critical role of human experts, at every stage of the GenAI testing lifecycle.

What comes next?

Pilot participants suggested 4 areas for future collaboration:

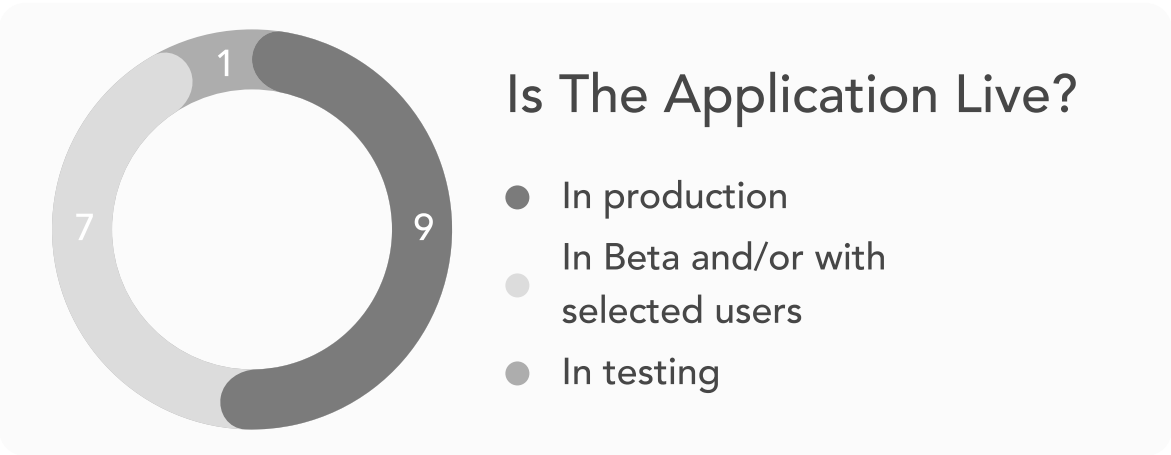
- Building awareness and sharing emerging best practices around GenAI testing
- Moving towards industry standards around "what to test" and "how to test"
- Creating an accreditation framework for testing
- Supporting greater automation for technical testing

The launch of IMDA Starter Kit – for consultation – is an initial step to address some of these requests.

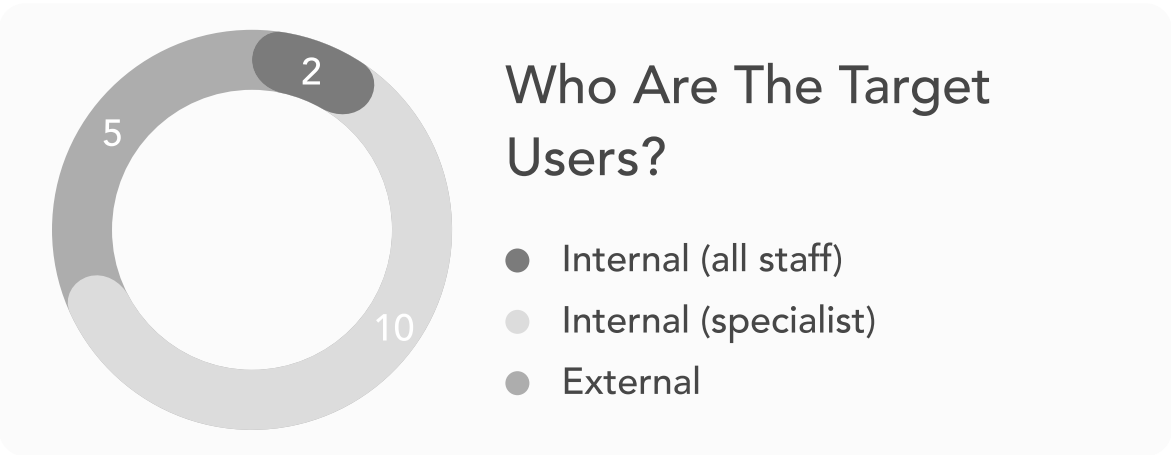
The journey towards making GenAI applications reliable in real-world settings has just started. IMDA and AIVF look forward to continued collaboration with AI builders, deployers and testers, and policy makers, on this important initiative.

2.2 Use cases

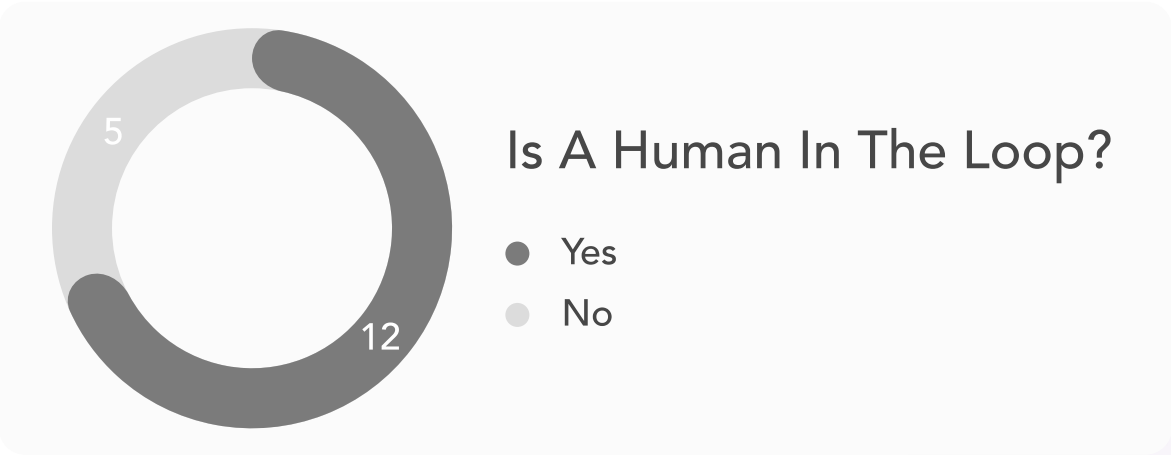
Background



16 of the 17 use cases were already live in production.
7 of them were in beta or/or rolled out to a limited group of users.



A majority were targeted at specialist users inside an organisation (e.g., software engineers at NCS). 5 were customer/ citizen-facing.



A human was “in the loop” in more than 2/3rd of the cases.
Even in the remaining 5, there was significant human involvement outside the immediate workflow of the application.

Full list of use cases

#	Tester(s)
---	-----------

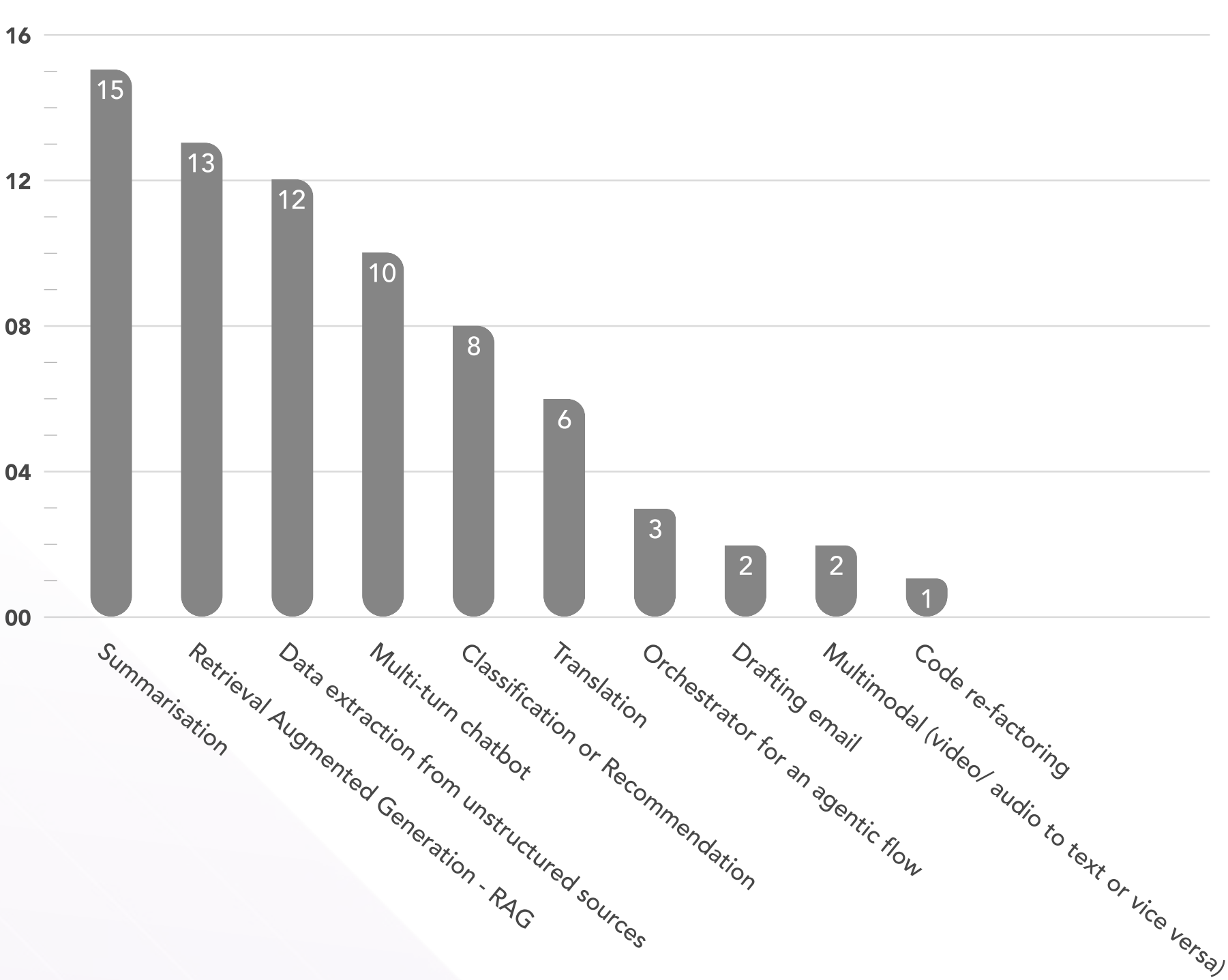
2.3 Patterns of LLM usage

Across the 17 applications, LLMs³ were used in diverse ways.

The top 3 usage patterns were Summarisation, Retrieval Augmented Generation and Data Extraction from unstructured sources. These patterns align with the focus of many of these applications on staff productivity improvement.

LLMs were also used to power multi-turn chatbots, and to help translate between languages. Relatively few used LLMs as part of agentic workflows – yet.

How are LLMs used in the application?



The table below maps each of the 17 applications to the different LLM usage modalities:

Table - How Are LLMs Used in the Applications?

Tester	Deployer	SUM	TRA	DAT	CLA	MUL	RAG	CHT	COD	AGE	EML
Advai	Checkmate										
AIDX	Fourtitude										

How did the pilot participants evaluate test results? (Number of Use Cases)

Tester	Deployer	Human judgement or review	Rule-based logic	Surface-level/ Semantic metrics	LLM-as-judge	Non-LLM model as judge
Advai	Checkmate					
AIDX	Fourtitude					
AIDX	Synapxe					
AIDX/Aiquris	ultra mAInds					
Fairly	Mind Interview					
Guardrails PRISM Eval	CAG					
Knovel	HTX					
LatticeFlow	European FI					
Parasoft	NCS					
PwC	SCB					
PwC	MNC Bank					
Quantpi	Unique					
Resaro	MSD					
Resaro	Tookitaki				*	
Softserve	CGH				*	
Verify AI	Insurer					^
Vulcan	High-tech Manufacturer					
Total		14	9	5	13	4

Legend

- * LLM used to extract facts as part of eval
- ^ Statistical models used to check effectiveness of LLM-as-judge

