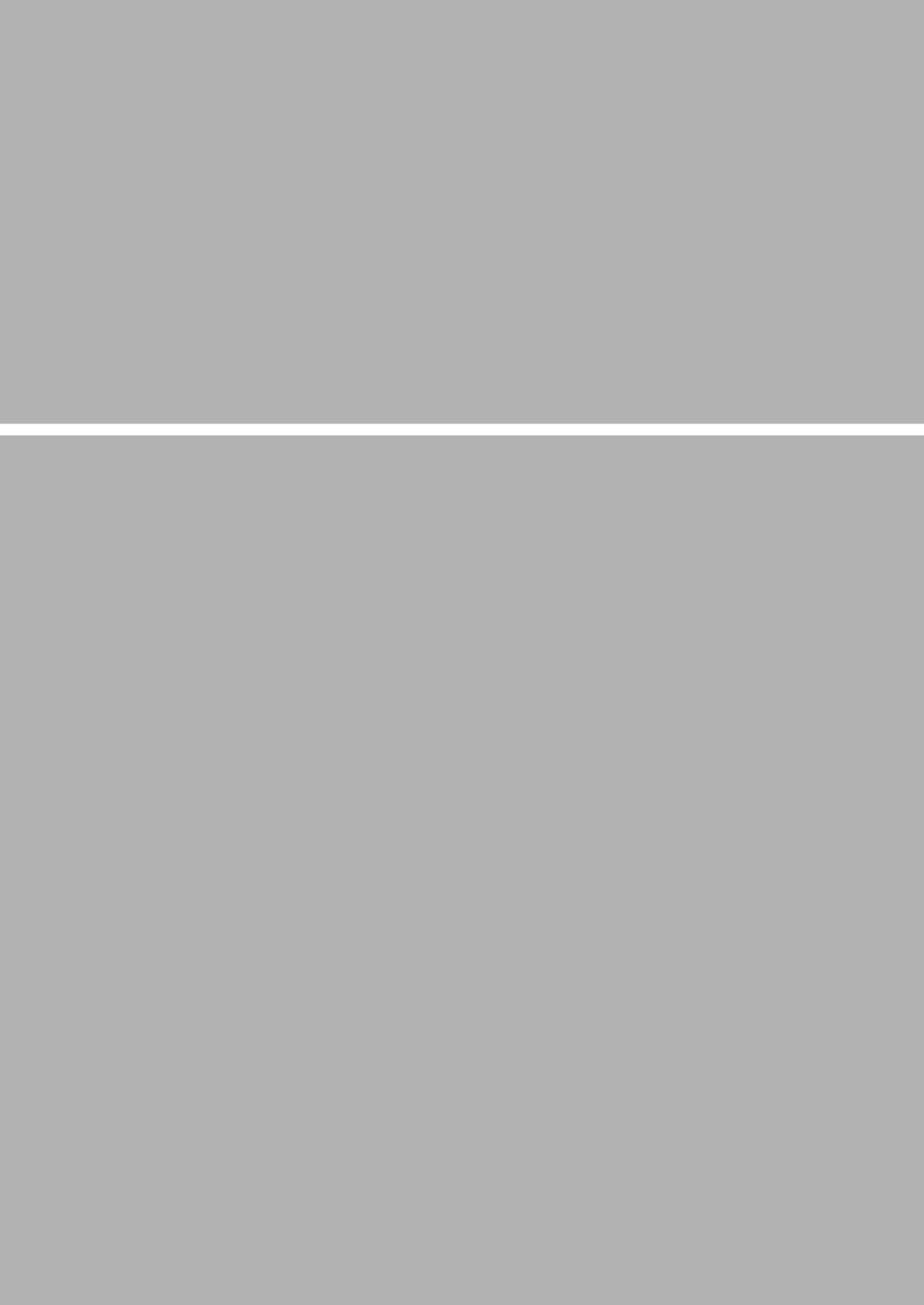




Chapter 4

The Future Will (Not) Be Annotated

INTRODUCTION

In February of 2023, the Singaporean government installed ChatGPT into the word processing software of 90,000 of its civil servants' computers. The aim of this trial was to speed up their research and writing work that would ease the increasingly large workload placed upon civil servants. As the manager of the scheme Mr. Soh put it, "We want to free officers up for higher-level tasks. This bot can help them get over that tough first draft or speed up their work by creating sample e-mails or even speeches⁶⁰." The following May, the Ministry of Communications and Information released guidelines for civil servants using ChatGPT to try and avoid the leaking of sensitive data or misuse of copyrighted material. The guidelines came with a reminder that they took personal responsibility for the accuracy of any text produced with a generative AI (gAI) like ChatGPT⁶¹. A year after the integration of ChatGPT into government software, it was reported that the Singaporean state, via its investment company Temasek, was in talks to invest in OpenAI, the creators of gAI. As Singapore attempts to become an AI hub in Southeast Asia, it seems that the future, and the future of its government operations in general, is increasingly tied to OpenAI and its flagship software.

Going forward, therefore, it is safe to assume that all government texts will be produced by or at least with gAI software from OpenAI. As shown in the previous chapter, AI systems are already being integrated into the Fatwa processes of MUIS built on ChatGPT software. Furthermore, speculation about how ministers' speeches are written has already begun. One interlocutor told me that he had been testing ministerial speeches with AI detection software and had already found several instances where they had been created wholly or in part by ChatGPT. This raises some significant challenges. As several reports have already demonstrated, gAI frequently produces text that recreates bias and negative stereotypes (Williams 2023⁶²). Therefore, when reading government policy and pronouncements, it will be prudent to consider not only the bias and ideology of the source, but also the software they used. This consideration now also applies to texts produced by MUIS given how its processes are becoming entangled with AI technology.

60 Chia, O. (2024). Civil servants to soon use CHATGPT to help with research, speech writing, The Straits Times. Available at: Accessed: 4 April 2025.

61 Chia, O. (2023). Public officers can use CHATGPT and similar AI, but must take responsibility for their work: MCI, The Straits Times. Available at: <https://www.straitstimes.com/tech/public-officers-allowed-to-use-chatgpt-and-other-ai-but-must-take-responsibility-for-work-mci> Accessed: 4 April 2025.

62 Williams, D. (2023). Bias optimizers. American Scientist. Accessed: 4 April 2025.

This presents a particularly fraught situation for Muslims in Singapore. This is because the Muslim community finds itself in the unique position of being a minority whose religious affairs are governed and regulated by the secular Singaporean state. For example, as discussed in the prior chapter, Friday sermons read in Singaporean mosques are vetted by the state. Two of my interlocutors, in fact, had already questioned whether Friday sermons were written by ChatGPT. This creates a situation in which Singaporean Muslims must ask themselves whether their religious proceedings have been written by a system, which has been shown to frequently generate Islamophobic jokes⁶³. As I will explore, this anxiety is only amplified by the existing anti-Muslim biases in Singaporean society.

During my time in Singapore, I wanted to explore how government texts, and the imagined futures they express, have changed with the introduction of ChatGPT. However, access to civil servants and their software is obviously difficult. While discussing these problems with my interlocutor Ahmad he made an interesting suggestion. Ahmad, who teaches English literature, argued that the best response would be to treat ChatGPT-generated texts like any other text throughout history. This does not always happen, as Cearns shows, where users of LLMs often perceive them as “a rational, scientific tool” (Cearns 2024, 69) with “intelligence [that] was ‘more than’ that of any one individual” (Cearns 2024, 66). However, as I have argued, AI is fallible in the text it produces, and therefore as Ahmad argues we should consider:

“...how we analyse it, what really is its standpoint, how do we feel about this, because feelings are something that ChatGPT cannot do right...my affinity to ChatGPT has always been, I’ll see what it can produce, and I will see how different that is from what I am expected to teach my students.”

In this chapter, I attempt to follow Ahmad’s suggestion. What would happen if I presented a ChatGPT-generated text that imagined the future of the Singaporean Muslim community to Singaporean Muslims themselves and asked them to annotate it? In this chapter, I will go on to explain how we can approach ChatGPT’s database as an archive that informs the texts it produces. What could taking this perspective offer us in a critical assessment of what ChatGPT writes? To do so I explore

63 Samuel, S. (2021) AI’s Islamophobia problem, Vox. Available at: Accessed: 4 April 2025.

the archival turn and various suggestions made to approach the archive not as neutral, but as a site where power can be exerted and contested. I then outline the possible use of annotation to expose the power relations implicit in the archive. I then present a selection of annotations made by my interlocutors.

In the following three sections, I explore the consequences of the observations made in these annotations and how they explain various aspects of gAI-assisted future-making in Singapore. The section fictions and speed explores how the pressure on civil servants to work faster can cause minor issues to result in major material consequences. The section “Why can’t the hijab be modern?” shows how ChatGPT repeats some of the most common tropes about Singaporean Muslims and how that the community trapped in the past. Finally, the section “Is there going to be mass conversion of Chinese people to Islam in the future?” demonstrates how ChatGPT’s mistakes are inextricably linked to the nature of texts that gAI attempts to emulate. I conclude by exploring what these issues may mean for the imagined futures of the Singaporean Muslim community. Firstly, though, I offer a brief outline of how the systems behind ChatGPT work.

ChatGPT: An Outline

On a basic level, ChatGPT is a relatively simple system. Based on the input data, or tokens, the user enters, the system predicts the next most likely word in the sequence based on its dataset (Radford et al. 2019). An example from Daniel Dugas, a researcher in machine learning and robotics, shows that if we input the statement “not all heroes...” into ChatGPT, then it will predict the most likely words to complete the sentence would be “wear capes”⁶⁴. This is why ChatGPT is often derisively referred to as a complex autocomplete. However, from the example above, we observe that ChatGPT does not really predict single words but instead a sequence of words. Each prediction is based on two factors: the likelihood that this word would follow the preceding word and the likelihood of those words to appear in its data on specific topics. In essence, the system focuses on certain words within the input data to create a probable output. For example, the prompt I used for this project was “What does the future of the Muslim community in Singapore look like? And can you describe a scene in that future?” In generating its response, ChatGPT may home

64 Dugas, D. (2020). How Deep is the Machine?, The GPT-3 Architecture, on a Napkin. Available at: https://dugas.ch/artificial_curiosity/GPT_architecture.html#paper1. Accessed: 4 April 2025.

in on “Muslim,” “Singapore,” and “future.” As a result, when predicting the sequence of words, ChatGPT will look at the data within its corpus relating to those topics.

It is how these systems make their predictions using Artificial Neural Nets (ANN) where things become more complicated. ANN are composed of layers of nodes that mimic a simplified form of human neurons. In a simple ANN, there can be as few as three layers: an input layer, a hidden layer, and an output layer. The input layer is the data, which is available to the system, and in the case of ChatGPT that is around 300 billion words⁶⁵. The hidden layer is a set of nodes which “allows the model to describe and predict complex relationships within the data” (Joque 2022, 46). It is in this hidden layer where the prediction happens. To take an example from Joque: if an ANN were being used to predict Pop-Tart sales after a hurricane, “one of the...neurons could be on the lookout for hurricanes so devastating that inhabitants leave town rather than stock up on food. We would want that neuron to detect when all three inputs...are very high and then suggest to the output neuron to decrease the predicted pop-tart sales” (Joque 2022, 47). As Joque points out, this is a simplification, and it is unlikely that a node would be looking for something so straightforward. In fact, this layer is referred to as hidden because there is little understanding of exactly what happens in this step. Instead, ANN are usually trained on questions with known answers, with the weighting of the hidden layer tweaked until it arrives at the correct outcomes. For example, the Pop-Tart ANN would be trained on historic data of Pop-Tart sales after a hurricane. This model is then applied to questions with unknown answers; this is known as supervised learning. ChatGPT also makes use of unsupervised learning, where there is no target answer, but the model is free to make any associations (Radford et al. 2019).

To grossly simplify, ChatGPT takes the input data from the user and subsequently focuses on a few key terms within that text to narrow its corpus of data. The ANN then makes complex predictions about the most likely sequence of words to appear in response to the user’s prompt. This sequence is then delivered to the user. This is problematic for a number of reasons, perhaps the most well-known of which is the tendency for this system to ‘hallucinate.’ Hallucination refers to situations where ChatGPT generates seemingly plausible text that is false. This happens as a result of the system generating probable sequences rather than factual

65 Hughes, Alex. 2023. ChatGPT: Everything you need to know about OpenAI's GPT-4 tool. BBC. . Accessed: 26 January 2026.

sequences. For example, when generating anthropology citations, ChatGPT may see that the most common name is Smith, and the most common year of publication is 2012. It would therefore produce a citation like “(Smith 2012)” despite these citations either not existing or not relating to the topic of a text. Emsley (2023) found several such hallucinations when asking ChatGPT to help generate a scientific paper. Whilst these are referred to as hallucinations and are benignly treated as glitches, they in fact reveal what ChatGPT is designed to do. Such citations have the correct form of citations, but they do not have the correct content or understanding of what citations are. As a result, what ChatGPT produces are simulacra of texts that are often correct in form but not content. ChatGPT is referencing everything and nothing at the same time. How might we engage with such texts? One approach could be to conceptualise ChatGPT as an archive and assume the critical methodologies used by archivists and historians.

ChatGPT as Archive

When OpenAI built GPT-3, they needed a mountain of data to do it. The system needed billions of example words and sentences in order to replicate human text. To be precise, the training dataset for GPT-3 had 499 billion tokens, roughly equivalent to 374 billion words. Sixty percent of those words came from an open-source database called Common Crawl, the data of which are collected by crawling webpages and scraping their content. However, OpenAI did not use this data raw; rather, it was filtered. In their explanatory paper on training ChatGPT, they explain that all the data from common crawl were evaluated for similarity to what they called “high quality reference corpora” (OpenAI 2020, 8). Most of these reference texts were drawn from English-language Wikipedia. In this paper, the process of constructing this dataset is presented as a neutral step in the larger process of building a Large Language Model (LLM).

I would argue that this process of assembling a training corpus is akin to archiving. As Carter states, the currency of the archive is “texts in the broadest sense of the term, including written, visual, audio visual and electronic” (Carter 2006, 216). Likewise, the authority of LLMs resides in the huge corpus of texts their algorithms replicate toward the production of objective policy. At one time the process of archiving was also seen as a neutral collection of texts. The archive was presented as a suppository of historical documents without political or power relations. Similarly, the database of ChatGPT is treated as a neutral collection of texts.

However, just as the texts within the colonial archive were procured via looting, LLM databases are produced via an analogous process of theft. For example, several AI companies are currently being sued for violation of copyright. This is not the only way in which LLMs and colonial archives are similar. As Arondekar writes about British India, “colonial governance was accomplished primarily through a practice of archiving, a systematic circulation, preservation, and recall of writing texts that allowed rule by remote control from London” (Arondekar 2009, 13). Archives under colonialism were a tool for knowing the other in order to govern them at a distance. As I have already argued in Chapter 2, AI and the Smart Nation within the Singaporean state operate similarly. LLMs and the colonial archive, then, are not just analogous of one another in that they share an epistemology. What might be gained, therefore, from viewing ChatGPT database as a kind of archive⁶⁶?

Since the archival turn, in which the role of archives was reassessed, the archive has ceased to be a neutral site of knowledge. As many historians, archivists, and anthropologists have reflected, the archive holds power in its ability to include and exclude. As Schwartz and Cook (2002) argue, the archive is a tool by which the powerful are able to enshrine certain narratives as ‘truth.’ As Foucault put it “The archive is first the law of what can be said, the system which governs the appearance of statements as unique events” (Foucault 1972, 129). For example, Stoler (2009) shows how by archiving colonial records, the uncertain social categories produced by colonial administrators became the rigid logics of empire. Once something enters the archive, it gains an aura of truth. In this way, inclusion is a way in which power is asserted in and through the archive.

Similarly, what is included in ChatGPT archive has a powerful impact on what this system produces. For example, as the OpenAI paper shows, “words such as violent, terrorism and terrorist co-occurred at a greater rate with Islam” (Brown et al., 2020, 38). This is likely a result of those words co-occurring with “Islam” at a high rate in the ChatGPT archive. This is significant because like how the archive gives text the aura of truth, text produced by AI systems is given an aura of objectivity. Joque (2022) calls this process, in which technology confers legitimacy, objectification. By using a machine, the subjective ideas of a text are made objects by being “carved into the material world” (Joque 2022, 82). In other words, because a machine produced a text, it is seen as more neutral or

66 There are some articles on using AI in archives (Jaillant, Mitchell, Ewoh-Opu and Urbaneja 2025), but nothing I could find discussed how AI can be viewed itself as a kind of archive.

objective. Further, just as the colonial archive informed imperial rule, we have seen how text generated from the ChatGPT archive is informing government policy.

However, the archive also has power through exclusion. As Carter (2004) points out, by excluding texts written by marginalised groups they can be silenced in the historical record. We could ask similarly pointed questions about what is included in the ChatGPT archive. Why is it that English-language Wikipedia is the mark of a quality text, and who is excluded when this is the filter? What other texts and perspectives exist offline that are not included in ChatGPT? In other words: what is excluded from the ChatGPT archive? The task Carter (2004) points out is in identifying these silences before these groups are completely forgotten from the archival record. How can such absences be identified? Carter argues that even when a group is entirely expunged from the archive, their absence can be felt in the silences left behind. Schwartz and Cook propose we do this by reading against the grain “to bring out voices which speak in opposition to power, or that insert irony or sarcasm or doubt” (Schwartz and Cook 2002, 15). In this way, they encourage the archive to be approached by seeking out traces of absence left behind.

How can we read against the grain of ChatGPT? It is not possible to access the database because this is proprietary material. Even if there were access to the database, it would be a huge task to explore what it contains. If someone had access, there could be a compelling ethnography of data scientists within OpenAI seen through the analytical lens of the archive. Such a project could explore the professional practices and assumptions of data scientists at OpenAI⁶⁷, as Schwartz and Cook (2002) suggest should happen with archivists. However, I did not have access to the database, or to OpenAI engineers. Instead, I attempted to read against the grain of the archive by inviting my interlocutors to annotate ChatGPT text. However, whilst such readings are usually argued to expose silences, and absences within the archive, I agree with Arondekar that implicit in such approaches “is the abiding assumptions that the archive...still constitutes *the* source of knowledge” (Arondekar 2009, 6). Subsequently, I take up Arondekar’s suggestion to view the archive as a system that produces fictions, which have been treated like facts. To do

67 This is, in fact, something that Nick Seaver (2022) proposes in his article on data-ethnography and in his book *Computing Taste* in which he spent time with programmers at Spotify to understand their approach to data. However, he does not approach these works through the analytical frame of the archive.

so, I ask my participants to annotate ChatGPT generated text in order to demonstrate the fictions AI systems produce.

Annotation is by all definitions a move by readers to engage with a text, whether that be marking it to remember (Jackson 2001), to correct the original text, (Orgel 2015) or to transmit information to other readers (Scott Warren 2010). Regardless of the annotator's intentions, all these acts serve to decentre the assumed authority of the printed word, instead drawing attention to the margins. Consequently, annotation has been used as a way of reading against the grain of the archive. For example, the artist Alexandra Bell (2018) has used annotation to correct *New York Times* articles about the murder of Michael Brown by police. By annotating these texts, Bell undermines the aura of objectivity which newspapers like the *New York Times* project, (Chai 2021) and exposes the history of racist texts, speaking to a kind of archive that implicitly informs the writing presented in the *New York Times*. By offering discursive space, annotations create an opportunity for people to speak back to archives and address their silences.

Just as the archival turn challenged the notion that archives are a source of objective truths, I hope that this chapter can demonstrate how ChatGPT, viewed through the lens of the archive, can be seen as reproducing power. This, I hope, will open a conversation about how to use ChatGPT more critically in the same way that many still use archives with a critical approach.

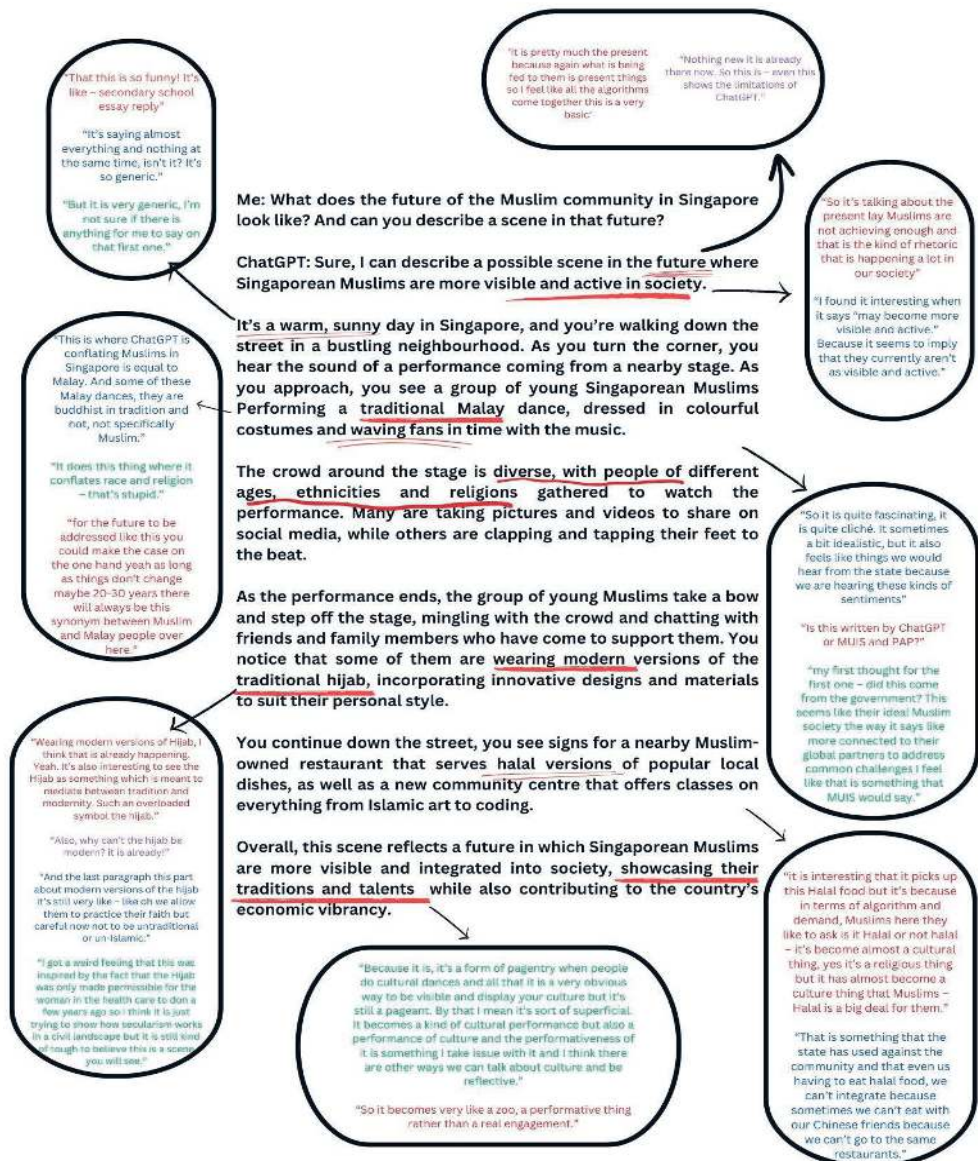
To this end, I prompted ChatGPT to predict the future of the Muslim community in Singapore. I then took this text with me to each of my interviews in Singapore and asked my interlocutors to annotate the text, showing the fictions that ChatGPT produces based on existing biases within its archive. These annotations open up the possibility to access the ChatGPT archive and understand the underlying assumptions embedded within its systems and sources. Further, by moving into this discursive space, my interlocutors are enabled to say, "we are here." Both these acts help to challenge the objectivity and finality of ChatGPT's pronouncements. As with archives before them, ChatGPT elevates itself beyond discourse by using technology to appear objective. As with archives, annotation can pull these gAI texts back into discussions about power. The annotations that follow aim to allow that discourse to occur so that Singaporean Muslims can correct mistakes in ChatGPT's record.

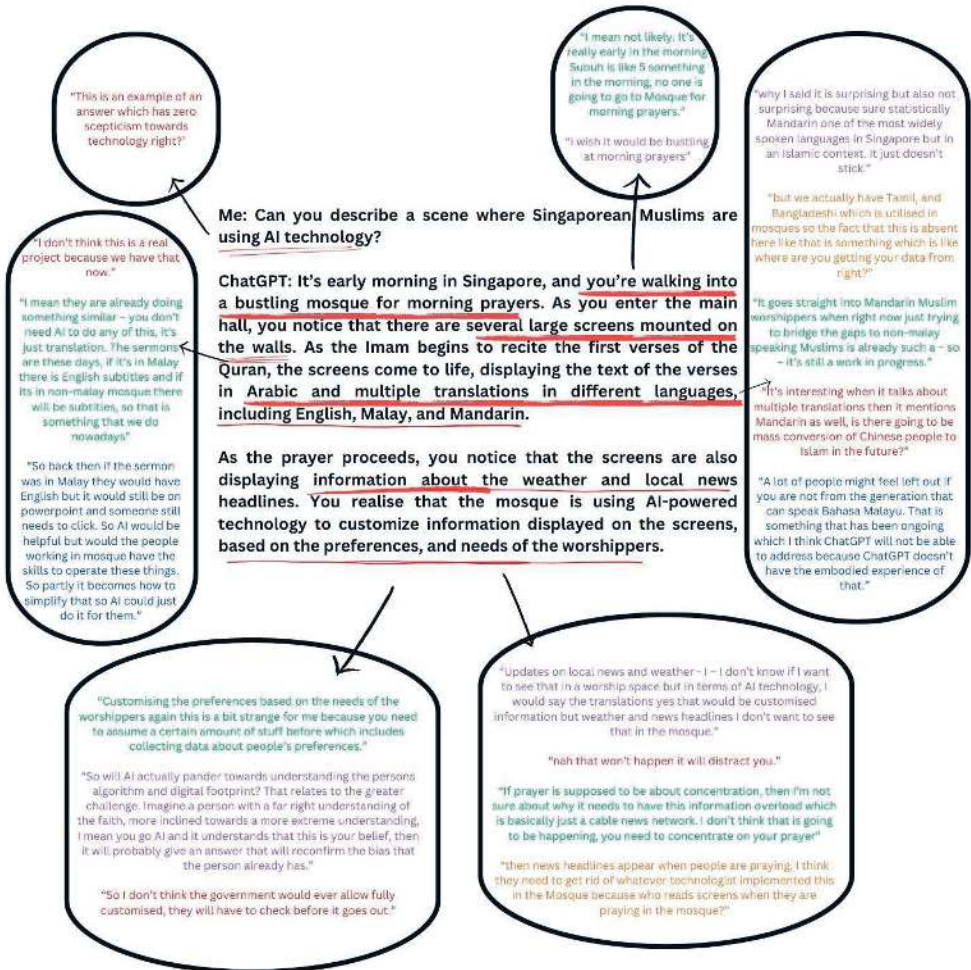
Further it presents them both as subject to the future ChatGPT imagines for them, as well as actors in the present and future.

During my fieldwork interviews, I would invite my interlocutors to read the ChatGPT-generated text, which is presented below. I prompted ChatGPT to imagine the future of the Muslim community in Singapore. Originally, I had planned to run a workshop for people to annotate and discuss the text together. However, due to various administrative reasons this was not possible⁶⁸. Instead, the annotation was done individually, and the annotations were later compiled into one text by me. Those annotating were from a mixture of backgrounds. Ten were Ustaz who worked with MUIS on a voluntary basis. Another ten were writers, activists, and artists. Five worked in the tech industry in various capacities. Finally, three university students annotated the text. In total, twenty-eight people read and annotated the text. I could not include all the comments here, but I tried to include comments which were representative of the sentiments shared.

The version of ChatGPT I used in March 2022 was 3.5. As of writing this chapter, the commercial and licensed model has since been updated with access to more data. Further, I prompted these responses at a time when there was little to no conversation around prompt engineering. This certainty limited the system's ability to produce the best possible results. However, for reasons I will return to, I believe these limitations do not prevent this exercise from giving us a glimpse into the AI archive.

⁶⁸ This being said, I would like to repeat this experiment with the latest version of ChatGPT and in a workshop setting in order to gain greater insights and develop the methodology further.





"I Mean Not Likely": Fictions and Speed

What do these annotations teach us about ChatGPT's archive? Perhaps most obvious to those familiar with Singapore, or Islam, are the small fictions. Throughout the text there are ideas that demonstrate this version of ChatGPT's misunderstanding of Islam and the Singapore context. These issues vary in seriousness. For example, one could argue that its claim of a "bustling mosque for Morning prayers" is an attempt to be literary. In reality, mosques are rarely 'bustling' for morning prayer because Fajr, the first prayer of the day, is usually around 5:30am. Consequently, it is far more likely that people perform these prayers at home. The response from my interlocutors was unsurprisingly light-hearted, with one saying, "I mean not likely," and an Ustaz telling me that he wished for such attendance at morning prayers. However, in this section

I aim to show how these seemingly minor mistakes can have a significant impact in the governance context of Singapore, where ChatGPT is increasingly becoming integrated.

Singaporean Muslims find themselves in a unique position. Singapore has an accommodative secular constitution (Article 15: 1; 4) and a majority non-Muslim Chinese population. However, since 1966 all Muslim affairs in Singapore have been managed under the Administration of Muslim Law Act (AMLA) (Mutalib 2008, 51). As previously discussed, the introduction of this law has centralised Islamic governance within the Singaporean state (Abdullah 2021, 20). As we already explored, the centralisation of Islamic religious governance allows the state to exert what Müller (2018) calls classificatory power in which the limits of Islamic identity are defined by state policy.

It is in this context that ChatGPT's 'errors' in this text become significant. In the section of the generated text where ChatGPT imagines a future where the Singaporean Muslim community uses AI technology, it reads: "You realise that the mosque is using AI-powered technology to customize information displayed on the screens, based on the preferences and needs of the worshippers." On the surface, my annotators, who are practicing Muslims, or themselves Ustaz, mocked this as a silly mistake like the bustling mosque, with one saying, "they need to get rid of whatever technologist implemented this in the Mosque." This annotator went on to explain that having screens within the mosque showing news did not make sense because this would distract worshippers from prayer.

However, several of my interlocutors disputed this on the grounds that the government would never allow such decentralisation. As one put it, "They would have to check it before it goes out." Indeed, this seems like an unlikely future, given that all sermons are checked and controlled by the government. As we have already explored, MUIS has, since 2003, introduced the Singapore Muslim Identity (SMI). The SMI comprises of ideals each Singaporean Muslim should embody in their practice. To try and encourage the SMI, MUIS employees told me that they use search engine optimisation practices to ensure Singaporean Muslims see MUIS rulings, in line with SMI, first when the latter google religious questions. Therefore, without tight government control, the personalised mosque screen technology ChatGPT describes is highly unlikely. This "error" represents the biases of ChatGPT's archive in two ways. First, the suggestion of TV

screens in the mosque suggests⁶⁹ that ChatGPT 3.5 does not have a large corpus on everyday Islamic religiosity, because if it did, it would perhaps notice that believers try to avoid distraction during prayer. Second, we can see that ChatGPT 3.5's archive does not have much information on a Singaporean context where religious messaging is tightly controlled; otherwise, it would not suggest personalised AI-generated messages. That context is important for Muslims in Singapore. After reviewing the ChatGPT-generated text for a while, one religious educator I met summed up his issue with it like this:

"If someone asks [AI] about celebrating⁷⁰ Christmas [as a Muslim], I think the answer would be you can't do it. But this is an issue that has a variety of opinions. By giving this one opinion can lead someone to think, 'This is the only opinion.' And if I'm asking this question...and the platform tells me I can't celebrate Christmas, it's hard for me to live in a multicultural society in Singapore."

This may seem like a strange example to give, but it refers to a specific well-known incident in Singapore. In 2015, the Zimbabwean preacher Mufti Menk was invited to speak in Singapore (Abdullah 2017). He was denied his permit on the grounds that he advised Muslims not to wish Christians "Merry Christmas." Opinions on this issue are divided (Abdullah 2021), and the previous Mufti of MUIS Sheikh Syed Isa Semait even wrote a Fatwa which aligns with Menk's opinion in the 1980s (Hassan and Bahari 2019⁷¹). However, the government stepped in during the Menk incident to emphasise that wishing others a Merry Christmas was imperative in order to uphold racial harmony in Singapore. Anything else would be seen as sowing the "seeds of terrorism and intolerance" (Ab Razak 2019, 423). This incident serves to show just how important context is in Singapore, where seemingly minor mistakes like advising against saying Merry Christmas can lead to significant consequences and even accusations of extremism.

69 I use this tentative language on purpose because we will never know for certain what is contained within the datasets of large language models. These annotations serve as an attempt to read against the grain of the ChatGPT archive in order to imagine what might be contained within its dataset.

70 Here my interlocutor mentions 'celebrating' Christmas. By this I interpret them to be referring to the issue of wishing others a merry Christmas and not literally to celebrate it as a religious holiday.

71 Hassan, M.H. and Bahari, M.B. 2019. A balance approach to the issue of 'merry Christmas' greeting by Muslims, Pergas Blog. Available at: . Accessed: 8 April 2025.

However, surely, if these errors can be spotted in quick annotations, then they will be easily detected and have limited impact. Further, whilst my interlocutors were all practicing Muslims, many with advanced degrees in theology, those operating the AI Fatwa bot described in the previous chapter will also have these qualifications. However, it is not only within MUIS that the Singaporean state makes decisions about the Muslim community. In some cases, it is less likely that Muslims will be in the room to correct the record. As one interlocutor pointed out to me, in their job at a large tech company, they often found themselves being the only Muslim in the room during policy conversations. They told me that being in these rooms is important because:

“Having that voice of Muslims [is important] because we sit in teams with so many languages and religions. You can’t expect everyone to know, but [they] also have their own biases. So, there have been instances where I do get questions like ‘is this normal?’”

Without Muslims in the room, or any other minority for that matter, it is easier to potentially misunderstand Muslim life. Given the integration of ChatGPT into everyday civil service operations, one annotator was even more cynical, arguing that ChatGPT would replace representation by being consulted in the way community leaders once were. Regardless, without fact checking, ChatGPT can propagate falsehoods which have tangible consequences.

Even with the proper representation, there is one trend that comes with all technology, ChatGPT included, which cannot be addressed: speed. The sociologist Hartmut Rosa argues that society is perpetually speeding up, a result of the need for “growth, innovation, and acceleration for its structural reproduction and to maintain its socioeconomic and institutional status quo” (Rosa 2016, 31–32). Technology, he argues, facilitates this speeding up in ways we all experience. For example, I am sure we have all felt the pressure to respond to email 24/7 because our devices make that possible. Likewise, ChatGPT offers “instant answers⁷²” (OpenAI 2022) that often sound plausible, and Singaporean civil servants are encouraged to speed up their research and policy writing by using the software. As a result, ChatGPT will facilitate the speed up of policy writing, which may in fact increase pressure on civil servants. As one of my interlocutors and a political scientist argued, “Knowing how civil

72 Open AI (2024) ChatGPT OpenAI.com. . Accessed: 8 March 2024.

servants are swamped with work...They may try to find shortcuts, and so I find this worrying.”

This worry is shared by Human-Computer Interactions scholars who argue that the most effective way to tackle online misinformation is to consciously slow down (Wilner et al. 2023; Fazio 2020; Salovich and Rapp 2021). ChatGPT used to offer a warning below its generated texts that read, “ChatGPT can make mistakes. Consider checking important information” (OpenAI 2023). However, as of 2025, this warning has been removed and replaced with a link to their terms of service in which they warn, “Given the probabilistic nature of machine learning, use of our Services may in some situations result in Output that does not accurately reflect real people, places, or facts⁷³.” These warnings are buried in the terms of service, so it seems unlikely that civil servants using ChatGPT will heed them when other societal pressures discourage slowing down and while these systems were introduced to enable an increase in speed. As a result, ChatGPT’s mistakes about the basics of Islamic faith and religious practice are bound to end up appearing in government documents, policies, and laws. Of course, the replication of racial biases within government documents is nothing new. It is just as likely that a civil servant in the Singaporean government does not know how busy the mosque is for Fajr prayers. What is different, however, is the way in which these errors can pass by undetected as government work becomes increasingly automated.

Further, these are not just innocuous mistakes about when people are likely to be praying in the mosque. They are also fundamental misunderstandings of how Islam is governed in Singapore, and the context which plays a central role in how that governance is managed. These ‘mistakes’ coupled with the pressure to speed up mean that ChatGPT’s unreliability stands to cause tangible negative outcomes for the Muslim community, given how entangled Singapore’s governance structures have become with OpenAI and its generative AI products.

“Why Can’t the Hijab Be Modern?”

In addition to the factual mistakes within ChatGPT’s text, there are also several instances in which it perpetuates negative stereotypes about the Muslim community. Chief among these is the repeated conflation of Muslim religious identity and Malay ethnicity. To quote one annotator,

73 Open AI (2024) *Terms of Use* OpenAI.com. Accessed: 8 April 2025.

“It does this thing where it conflates race and religion – that’s stupid.” Where my prompt requested a scene depicting the future of the Muslim community in Singapore, ChatGPT’s text describes the performance of a Malay dance. This is problematic for two interconnected reasons. Firstly, this constant conflation of Malay-ness and Islam excludes large segments of the Muslim community, such as the Indian Muslim community who constituted around 69,000 people as of 2021 (Census, X 2021⁷⁴). Secondly, the denial of non-Malay Muslims’ existence expunges the unique concerns of these communities. This is especially egregious in the context of Singapore, where race and religion are central to navigating citizenship (Tschacher 2018, 17). This conflation results in the flattening of Muslim concerns while prioritising those of the Malay community. At best this represents a failure to understand the complexity of the community, and at worst it renders segments of the community invisible.

This flattening of Muslim identity appears in the way ChatGPT handles two topics: halal food and the hijab. In ChatGPT’s constructed scene, we find restaurants serving “halal versions of popular local dishes” and women “wearing modern versions of traditional hijab.” As annotators pointed out, halal food and the hijab have frequently been used as examples by the state as Muslim failure to integrate into Singapore. It has often been argued that Muslims are unwilling to eat with non-Muslims because of halal food (Musalib 2012). Similarly, the hijab has become a visual symbol of the Muslim communities perceived failure to integrate (Abdullah 2016). As the former prime minister Lee Kuan Yew argued in 2011, “I would say, today, we can integrate all religions and races except Islam” (Rahim 2012, 170).

By tying it to religious prescriptions, the ‘failure’ of the community to integrate is made into an intractable issue. Not only is the Muslim community not currently integrated, but also it does not have a route toward integration. As scholars like Barr (2013) argue, a vicious cycle begins where a perceived failure to integrate denies access to resources for the Muslim community. As a result, the Muslim community will “always be kept separate by various systemic programmes concerning language, public culture, education, and outright discrimination” (Barr 2013, 113). By suggesting that Singaporean Muslims have failed to integrate properly, ChatGPT’s text maintains the myth that Islamic faith prevents proper engagement in modern society.

74 SingStat (2020) Singapore census of population 2020, KEY INDICATORS OF RESIDENT POPULATION. Available at: . Accessed: 8 April 2025.

We can see this dominant narrative in ChatGPT's text about the future of Muslims in Singapore. Throughout the text, ChatGPT emphasises a positive, harmonious future in Singapore. However, this future is not achieved by addressing racial inequalities which impact the Muslim community but by Muslims becoming "more visible and integrated into society" (MUIS 2006). As one of my interlocutors commented, "It feels like it is very nicely painted, so I am very curious if you feed them more about the different nuances of the Muslims in Singapore, would they be [so nicely painted]?" By "nicely painted" here, they are not referring to the style but to the overwhelming positivity of the ChatGPT-generated text that neglects the material challenges faced by the Muslim community in Singapore. Following the official state line, ChatGPT regurgitates the idea that racial harmony is already achieved in Singapore and can only be undermined by the failure of Muslims to correctly integrate.

The conception that the Muslim community does not integrate well has its roots in the history of Singapore and the region more generally. Given the presence of Singapore's much larger Muslim-majority neighbours, Malaysia and Indonesia, there has long been anxiety about the loyalty of Singaporean Muslims. This anxiety is tied to the shared ethnic identity between the majority of Singaporean Muslims, 80% of whom are Malay, and Malaysian Muslims. As we have seen, the Muslim community in Singapore is perceived as a dangerous element that could cause a return to pre-independence racial violence. In this context, certain behaviours, such as wearing the hijab or not eating with non-Muslims because of concerns about halal, sow the seeds of division and are routes to violence. It is this suspicion that drives the pressure for Muslims to integrate as modern Singaporeans and simultaneously denies that integration.

However, as one annotator pointed out "why can't the hijab be modern? It already is!" It is a running joke that Singapore remains two to ten years from completion. The 'future' city is just out of reach. As Collins (2008) has argued about consumerist societies, the future is the products you do not own yet, and there is always a new product. As a result, we remain in a kind of perpetual present. This is true for Singapore, which is frequently on the move in search of the future. Likewise, Muslims in Singapore find themselves trapped in the present. The future, as presented by the state and ChatGPT, is one where Muslims are "more visible and integrated into society" (MUIS 2006) However, this future remains just out of reach because of the association between their religious identity and histories of racial discord in Singapore.

Therefore, there is no hijab sufficiently modern that can overcome the idea that the hijab represents a failure to integrate. As I have already shown, Bear (2016) calls these narratives epistemes of future in which apocalyptic futures are anticipated and against which states guard themselves. This focus on anticipating disaster creates a “continually receding horizon of the future [that] determines our actions in the present” (Bear 2016, 493). In Singapore, the apocalyptic future is one in which the delicate balance of racial harmony fails, leading to violence between races. Here the past, present, and future are collapsing into one moment, freezing the Muslim community in what Bear calls a “foreclosed future” (Bear 2016, 496) as described elsewhere in this thesis.

These tired tropes, and their debunking, go back to the time of Singaporean independence if not earlier. There are many articles covering all of the issues described in this chapter in much greater depth⁷⁵. The question, therefore, is what role ChatGPT plays in continuing to reproduce these negative stereotypes. As we have explored, ChatGPT fundamentally works on the basis of statistical probability. Based on its dataset, it produces text that reflects the most likely response to a prompt. This suggests that while many thousands of words have been spent debunking the stereotypes about Singaporean Muslims, much more data exists within ChatGPT’s dataset that confirms them. What matters most to ChatGPT is what things are most common rather than what is most accurate. Observing this led one annotator to refer to what ChatGPT does as regurgitation saying,

“I’m not sure there is anything for me to say about it. Largely because it is a very generic statement which is being made, and again, I am not even sure what kind of data is being fed in order to churn out these responses.... It’s about more than a thing being able to churn out or regurgitate⁷⁶ the things it has read—it needs to be able to read through and interpret, refer, and most importantly debate.”

ChatGPT can only make text that replicates and recombines its archive; therefore, it can only represent the present. This led to the most frequent annotation—that the future being presented already exists. As a result,

75 Namely, Abdullah 2016, 2017, 2021, Alatas 1977, Chua 1995, 2003, Febrica 2010, Ganapathy and Fee 2016, Müller 2018, Mutalib 2008, Noor 2011, Rahim 2012, Rahim 2020, Rashid 2016, Seng 2009, Tan 2016, Zainal and Wong 2017, to name just a few.

76 It is from this quote that I draw the inspiration for my thesis subtitle, *AI Regurgitation*.

ChatGPT is bound to replicate the foreclosed futures Bear describes, in which Muslims are never quite modern or integrated enough.

“Is There Going to be Mass Conversion of Chinese People to Islam in the Future?": Foundational Issues

So far, most of the critique I have presented of this ChatGPT-generated text could be dismissed on technical grounds. The version of ChatGPT I used was 3.5 from March 2022, the commercial and licensed model has since been updated with access to more data. Further, I prompted these responses at a time when there was little to no conversation around prompt engineering. Certainly, prompt engineering, where the user refines their question to improve ChatGPT outputs, could have increased the quality of the generated text. The argument goes that improving the technology and the skill of the user will address any and all of the issues raised by the annotations I have presented.

However, there is research to suggest that increasing amounts of data and ChatGPT updates have not improved the quality of its generated text. One paper found that, in fact, ChatGPT performance on certain tasks had worsened over time (Chen, Zaharia and Zou 2023). The *MIT Technology Review* has reported since late 2022 that OpenAI may be running out of quality data to train their model on (Xu 2022). Furthermore, there is evidence that using AI generated text to plug the gap left by man-made text, may lead to model collapse⁷⁷ (Shumailov et al. 2023). Most damning for the argument that systems like ChatGPT will inevitably improve is a study from researchers at Radboud University, which shows that AGI (Artificial General Intelligence), OpenAI's stated aim, is computationally intractable (van Rooij et al. 2023). This means that the quantity of data and computational power required to achieve “sentience” through systems like ChatGPT are unachievable. These articles taken together suggest that we are nearing the ceiling of what LLMs such as ChatGPT can achieve. However, in this section I want to argue that ChatGPT texts have more foundational issues when it comes to generating novel text, which cannot be addressed by technical improvements.

77 Model collapse as described by Shumailov et al (2023) is where LLMs slowly degrade over time because the data input is not high quality. Over time, bad data erodes the ability of the LLM to write convincing human text. Eventually, this would theoretically lead LLMs to produce nonsense sentences with no grammatical or lexical logic. This is particularly interesting in my case because Shumailov et al (2023). show that model collapse appears first in minority data sets on topics which have little or no data within the ChatGPT archive. Consequently, LLMs tend to produce hallucinations and mistakes first about minority groups, like Muslims in Singapore, which are less often noticed because the model as a whole appears to be performing well.

Whenever someone writes a text or creates an archive, they make choices about which information to include and exclude. There is no text which is exhaustive, and what appears in a text is always a choice. This becomes an issue with generative AI text because readers do not have access to the reasoning used by ChatGPT when making these choices. As Salvaggio aptly puts it, generative AI outputs “are infographics that don’t tell us how to read them...imagine a map that doesn’t have North, South, East and West. We have to work this out by looking for certain landmarks that we do understand⁷⁸” (Salvaggio 2023). ChatGPT-generated texts are representations of the dataset they are built upon, which is in turn is an incomplete sample of the real world. This is another way in which ChatGPT is similar to archives, which Verne Harris has argued, “is best understood as a sliver of a sliver of a sliver of a window into process” (Harris 2002, 84). In the annotations of ChatGPT presented here, then, we are searching for landmarks that might give us a window on to this sliver of a sliver of a process.

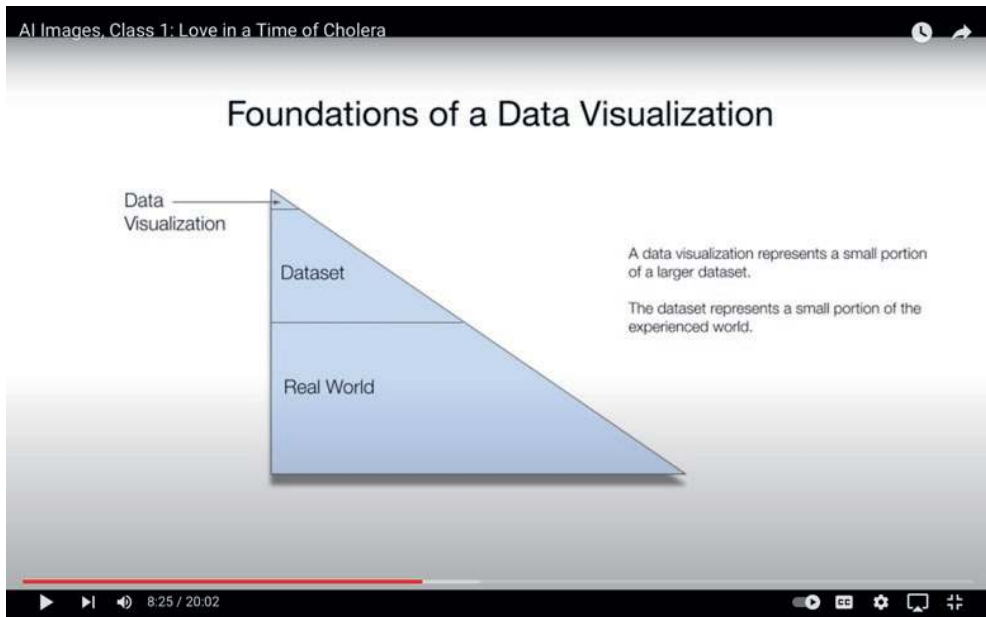


Figure 10–The AI database iceberg

One landmark that may help us understand the choices ChatGPT made in this generated text are the three languages it mentions being translated in the mosque, namely Malay, Mandarin, and English. As one of my in-

⁷⁸ Salvaggio, Eryk. 2024. AI Images, Class 1: Love in a Time of Cholera. YouTube. . Accessed: 26 January 2026.

terlocutors pointed out, “Sure, statistically, Mandarin is one of the most widely spoken languages in Singapore, but in an Islamic context, it just doesn’t stick.” Another jokingly said, “is there going to be mass conversion of Chinese people to Islam in the future?” Both are right to point out that it is highly unlikely that mosques would choose to translate to Mandarin, a language hardly spoken by the Singaporean Muslim community. However, if you google what the most widely spoken languages in Singapore are, you will find the answer: English, Mandarin, Malay, and Tamil. This suggests that ChatGPT chooses to present the most dominant or common narratives in its texts. This is further supported by the regurgitated stereotypes we already discussed. Due north for ChatGPT, then, means representing the majority.

What the North Star of the Mandarin translation in the mosque demonstrates is what our understanding of how ChatGPT already suggests. As the AI research group Estampa argue, “generative AI tools disassemble language (visual, textual) and reassemble it based on the calculation of probabilities” (Estampa 2024). What matters to ChatGPT, then, is not what is true but what is most likely. Striphas refers to this as algorithmic culture, meaning “the use of computational processes to sort, classify and hierarchise people, places, objects and ideas” (Striphas 2015). This results in the exclusion of minority perspectives because they are not well represented ChatGPT’s outputs. Even without a detailed understanding of the inner workings of ChatGPT, annotators point to several of these instances within this short text, including the exclusion of Indian Muslims whose language translation needs are not mentioned. This preference towards the most probable also has the interesting side effect of producing predictable and clichéd writing. This is something that is frequently commented on by annotators who derided the writing quality. The result is boring writing that fails to represent complex realities.

Of course, it does not need to be the case that big data sets exclude minority perspectives. As Foucault-Welles (2014) argues, large data sets can make otherwise difficult to access minorities visible. For example, she points to her research about women over fifty who have played online games for over 1000 hours. From an original dataset of fifteen million users, Foucault Welles was able to identify one thousand five hundred users who matched her research criteria, a group who usually “get lost in Big Data analytics, wiped away as noise among the statistically average masses” (Foucault Welles 2014, 2). However, as I have shown in previous chapters, and as the annotations of my interlocutors attest, this is not

the approach taken by LLMs like ChatGPT⁷⁹. These systems tend towards a conception of truth in which bigger is better (boyd and Crawford 2012). The consequence is that, in general, big data averages that exclude outliers are the norm in AI projects like ChatGPT.

There are already examples of this preference for bigger datasets in the form of Google Translate. In 2016, Google announced that it would begin using machine learning to improve the quality of its translations (King 2019). This change entailed a shift towards the kind of technology which underlies ChatGPT, in which the system would scan its corpus for the most probable translation of a passage. While this change improved the quality of translation in general, it led to some interesting mistakes in translation of literary texts and figurative writing. As Large (2017) observed when translating *Hamlet*'s "To be or not to be" monologue to German, Google Translate mistakenly turns the word "fortune" into the financial sense rather than fate. He concludes this is a result of Google Translate's corpus preferencing business and political texts. Critics argue that this change would lead to the erasure of more idiosyncratic language and idioms, which are not well represented within text, reducing the chances of surprising novel choices.

However, as King (2019) points out, Google aimed to guard against this outcome by encouraging users to submit alterations and annotations. The algorithm would make the translations broadly statistically accurate while human users return the specific and distinctive. A similar system could be introduced to ChatGPT through text annotation by users. It would be easy to dismiss this as idealistic; however, we have seen how "participatory web cultures" (Beer 2009) in projects like Wikipedia have found success. A somewhat similar process occurs with ChatGPT wherein

79 Further, increased visibility can be a double-edged sword. As Foucault Welles argues, the historic lack of data about women and minority groups has resulted in them being "problematically written into the margins of social science, discussed only in terms of their differences, or else excluded altogether" (Foucault Wells 2014, 2). For example, as D'Ignazio and Klein (2020) show, up until 2014 no one was tracking data on the mortality rate of black women during childbirth in America. This is in spite of several civil society groups campaigning about black maternal health for decades. However, increased diversity within datasets often leads to "more often than not...enlisted in the service of increased oppression, greater surveillance, and targeted violence" (D'Ignazio and Klein 2020, 14). For example, the Chinese startup CloudWalk aimed to produce a facial recognition dataset that did not exclude the faces of people of African descent. To do so, they signed a deal with the Zimbabwean government to provide facial scans, and in return the Government of Zimbabwe received "a national facial database and 'smart' surveillance infrastructure" (D'Ignazio and Klein 2020, 13). This presents an infrastructure which can then be used by the government to track citizens, a government which D'Ignazio and Klein note has a "dismal record of human rights" (D'Ignazio and Klein 2020, 13). Being more visible can also result in further repression, a fact which I return to in chapter 5 of this thesis.

user data are used in the feedback loop that improves ChatGPT's outputs. It could be argued, therefore, that there is a participatory aspect to LLMs like ChatGPT.

However, there are caveats to this utopian vision of an annotated ChatGPT. First, it is unclear how effective user edits of translations have been. While Constantine (2019) reports being impressed by Google's improved ability to translate Voltaire, other translators like King are less convinced: "I used the GT feedback tool to insert my English translation as the "correct" one. My feedback appears to have had no impact on how GT translated this text over 10 months" (King 2019, 81). Further, these tools are unpaid, a fact which allows Google to profit off the impression that their translation software is autonomous when, in fact, significant unpaid human labour is required for it to operate. This issue of invisible labour is also prominent in ChatGPT. As Estampa argues, "behind the autonomous appearance of AI tools...we find different levels of human resources displaced to different geographies, precured and invisible by the technology innovation industry" (Estampa 2024). Gray and Suri (2019) refer to this labour as ghost work, in which people are paid to perform tasks to enable machines to become automated. Whilst Gray, and Suri (2019) reflect on the advantages of flexibility such work offers, and that such workers find ways to collectively organise and survive the systems isolation, such systems often produce bad working conditions. For example, Gary and Suri (2019) argue that opaque working relations mean that ghost workers are frequently not paid for work they performed and had no recourse when this happened. By suggesting OpenAI integrate annotation into ChatGPT, we open the door for more of this kind of exploited, invisibilised labour within generative AI.

This opacity about labour hints at another issue: the general uncertainty around how these systems work. As King points out in her conclusion, "Beyond the debate over quality lies the question of transparency and intent: who is developing these ubiquitous tools? And how are they being used and toward what ends?" (King 2019, 88). We can see this opacity in the case of Constantine's (2019) study, where he attributes the improvement of Google Translate to the algorithm and not to the contributions of human users who more likely drove the changes. The same applies to ChatGPT. Due to the opacity of its operation, it is impossible to know which discourses its texts rely on and what its algorithm values most. What is clear, as the annotations of ChatGPT's texts reveal, is that these systems prefer the most common majority positions, or as Joque puts it,

that they value, “decisions the market will most likely reward” (Joque 2022, 32). However, we should be careful with such decisions that not only disregard minority groups who are seen as statistically insignificant but that also regurgitate negative stereotypes about them when they are represented. As Beer concludes, this process of “sinking of software into our mundane routines...means that these new vital and intelligent power structures are on the inside of our everyday lives” (Beer 2009, 993). In the case of Singapore, this software has sunk into the day-to-day governance apparatus.

It appears that ChatGPT has become increasingly integrated into the governing systems of Singapore⁸⁰. Of course, a limitation of my research is a lack of access to state processes. From the small window into ChatGPT provided by my interlocutors and their annotations, it appears that these systems and the civil servants within them are under increasing pressure to work quickly—something ChatGPT further facilitates. It could be the case that civil servants are not speeding up, as Rosa (2016) suggests they might be, but instead slowing down in order to annotate and critique ChatGPT outputs. However, as we saw in the previous chapter, MUIS intends to integrate ChatGPT into Fatwa decision making processes partially to speed up the work rate of the Fatwa committee. The intention, in part, is to speed up government processes—not slow them down.

In addition, in this section I have shown how creating texts from the archive necessarily replicates the information stored within that archive. This is equally true for ChatGPT, which tends to exclude minority views from its texts, preferring to reproduce the most common views. By doing so, ChatGPT replicates the foreclosed futures imagined for the Muslim community by the Singaporean state. By mechanising these processes, they can be treated as “objective” (Joque 2022). However, what is lost is human input on how these choices are made⁸¹. The argument, therefore, that improved technology and training users are enough to ensure their unproblematic use, is flawed. When these texts are used without critique, they are bound to perpetuate foreclosed futures. The resulting imagined

80 Chia, O. (2023) Public officers can use CHATGPT and similar AI, but must take responsibility for their work: MCI, The Straits Times. Available at: . Accessed: 04 April 2025.

81 This is another limitation of my research. At the time of my fieldwork, my interlocutors were not extensively using ChatGPT, and therefore it was hard to assess how they were using it. Furthermore, much more is now known about how to improve prompting. It could be the case that people are using AI systems in much more considerate ways, and further fieldwork would be required to explore this question.

futures are also harder to challenge because, as Salvaggio shows, we have no keys or legends to understand how they were written.

CONCLUSION: CHATGPT AND THE SINGAPOREAN MUSLIM COMMUNITY

After it became clear that ChatGPT 3.5 was spouting bigoted views⁸², including Islamophobic jokes⁸³, OpenAI applied content filters to their chatbot. Ostensibly, this solved the problem. ChatGPT no longer espouses its racist bias, at least publicly. Several commenters online refer to this new, nicer, ChatGPT 3.5 as a Shoggoth with a smiley face. A reference to Lovecraft's Shoggoth, an unimaginable cosmic horror that devours and destroys everything in its path except with the addition of a comforting emoji. When I first generated the text analysed in this chapter, it seemed benign enough, and some of my interlocutors even agreed. However, the longer my annotators looked, and as more people added their feedback, the roiling horror beneath the surface became clear.

What reading against the grain of ChatGPT's archive seems to suggest is that the LLM's database confirms the biased ideas of the Singaporean state. What we can say for sure is that these ideas are regurgitated by ChatGPT within its texts, and this suggests that much of its archive is made up of texts which contain such biases. Some may argue that more sophisticated filters, greater access to data, or numbers of parameters will solve this problem. ChatGPT's ability to produce unbiased, objective text is just one patch away. However, by treating ChatGPT outputs as text produced from an archive, this chapter aims to show how ChatGPT tends to reproduce power. Just as one more source within the archive cannot produce completeness or objectivity, one more piece of data within ChatGPT's database will not either.

82 Perrigo, B. (2023) OpenAI used Kenyan workers on less than \$2 per hour: Exclusive, Time. Available at: Accessed: 8 April 2025.

83 Samuel, S. (2021) AI's Islamophobia problem, Vox. Available at: Accessed: 4 April 2025.

gAI's inability to account for outliers means that it will continue to reproduce biased ideas common within its dataset. Its reliance on the most common existing data, means that it will be bound to regurgitating the stereotypes held by the majority against the minority. Finally, and most importantly, texts and their archives are a compromise in which the writer leaves out or includes information on the basis of their ideology. To return to a Derrida quote, text is "one discourse propped on another" (Derrida 1991, 197). In this way, ChatGPT texts are in continuity with archives and their perceived objectivity that preceded them. What ChatGPT lacks is accountability for those choices because they are often made by a machine and given the veneer of objectivity. To avoid these pitfalls, it is necessary for there to be a reconsideration of LLMs similar to the one that occurred in the archival turn. However, such an approach requires time, something that goes against the aims of AI systems like ChatGPT.

As we have seen, MUIS and the Singaporean state are committed to using generative AI. What will this mean for the imagined future of the Muslim community? At events throughout the city, I heard the phrase "AI won't replace you, but someone who knows how to use it might." Mostly, these critiques came from speakers at tech conferences and upskilling events selling the virtues of a digital future. However, for my annotators this "future" looks uncannily like the past and present. The words that came up most often, were "cliché" and "generic". These were platitudes they had all heard before, and beneath those were the echoes of racialised bias common in Singapore. As one annotator put it:

"For the future to be addressed like this [pause] you could make the case that as long as things don't change, maybe in 20 to 30 years there will always be this conflation between Muslim and Malay people."

Archives have always entailed inclusion and exclusion of ideas—decisions that are implicitly or explicitly driven by ideology. What ChatGPT changes is how it limits our access to understanding how those choices are made. Decisions about what information is valuable move into the hands of large tech companies that operate based on algorithmic cultures that value what is most common. Frequently, this allows mistakes, stereotypes, and simplifications like the conflation of Malay and Muslim in Singapore to proliferate. At the same time, because of its basis in probability, ChatGPT is endowed with the objectivity associated with mathematical certainty. As people are encouraged to work faster, and ChatGPT

facilitates that speed, more of these mistakes and biases will slip by unchallenged. Each time they do, the horizon of the future draws closer (Bear 2016). Without a critical engagement with the archive of LLMs, we risk the future becoming a perpetual regurgitation of the past, including all its mistakes. Perhaps annotation can be a tool to try and break open the black box of generative AI and allow us to understand more clearly how it reaches its pronouncements.

THE FUTURE, AN INTERLUDE: PART 2, RACE AND FORECLOSURE

The Singaporean state-owned electronic payment app E-pay released an advert in 2019 which included a Singaporean-Chinese actor, Dennis Chew, appearing as the four races which comprise Singapore's CMIO racial designation. When appearing as the Indian character, Chew's skin was darkened with make-up. When Chew appeared as a Malay character, he donned a tudung. In response, Indian-Singaporean rappers Subhas Nair and Preeti Nair made a parody song called "K. Muthusamy" (the name of the Indian character in the advert), with the chorus hook, "Chinese people always out here fucking it up." In response, both Subhas and Preeti were questioned by police and given a two-year conditional warning for inciting racial hatred. Later, in 2021, Subhas tweeted comments which were seen as violating that two-year warning, and he was sentenced to six weeks in prison, which he served in February and March 2025.

Then Home Affairs and Law Minister K. Shanmugam addressed the press in response to the original rap video, stating,

"Why take it so seriously? One video is not going to lead to violence... think of it like this, if we allow this, we have to allow other videos... what will happen to our racial harmony... suppose you allow this video attacking Indians, Malays... would minorities feel safe? ... There is a reason we have racial harmony in Singapore... We must maintain that, we will maintain that... we cannot allow these attacks⁸⁴."

At the 2023 AAA conference, Joshua Babcock analysed these comments, arguing that they represented an example of Singapore's anxieties about racial violence. He argued that, by making this slippery slope argument, Shanmugam was not only suggesting that violence was the inevitable outcome of videos which discussed race, but in fact collapsed the distance between discussing race and racial violence. Babcock focuses on Shanmugam's use of the word "attacks" throughout his speech, a word which suggests the violence has already happened.

84 CNA "Videos that 'attack another race' cross the line, says Shanmugam on Preeti's rap video". Accessed: 6 November 2025.

This incident, and the government response, I would argue, reflects what Crapanzano (2007) calls the thinning of the near future, in which people and governments become more concerned with the present and the distant future. Crapanzano (2007) points to America as an example of this presentism, in which some short-term threats are over-emphasised, such as terrorism, while long-term threats like climate change are denied.

In the controversy around the E-pay advert, we can see this short-termism in action. As Babcock argues, the anxieties around racial violence drive the response of the state, collapsing the distance between an incendiary comment and literal violence. The ironic outcome is that many feel race is not sufficiently discussed, allowing racial biases to persist and proliferate. Consequently, the near future, in which racial disharmony grows, is ignored and denied. In fact, Shanmugam would do exactly this in an interview days after the one discussed above, where he argued racism was reducing in Singapore. Meanwhile, the short-term danger of violence is over-emphasised.

As the previous chapter demonstrates, AI systems may in fact worsen this short-termism. As I argued, the racial categorisation of Singaporean Muslims within ChatGPT's dataset can be deduced from its outputs, which replicate ideas such as the notion that Singaporean Muslims are not "modern." ChatGPT is therefore regurgitating the ideas of race put forward by the Singaporean state. As we have seen, such systems are being integrated into government processes, both within religious institutions and in other ministries. Consequently, such ideas about race are likely to be codified and entrenched within these organisations.

This is not only a product of state ideology, but, as I argued in my chapter on the epistemology of AI, is a product of the way knowledge is being created within the state. As digital culture researcher Ramon Amaro argues, "'algorithmic bias' is based on the fundamental principle of mathematics... which is already predicated on the reduction of chance and contingency: namely, to regress, to clean, to normalise the unexpected, and therefore normalise and exclude the opportunities that live outside the lines of political engagement" (Amaro 2020, 155). In other words, in order to make predictions, the identities of people must be cleaned and made to fit within prescribed categories.

Again, we see here Bear's foreclosed futures. The possibilities of the future are reduced, both in the sense that violence is made inevitable, and that identity is proscribed. The limits of what it can mean to be Muslim in Singapore, and what futures can be expected, are closed down. Up until now, this thesis has focused on how the epistemology, infrastructures and ideologies that come with algorithmic governance foreclose the future. However, as I stated in "Future Part One", the future is still open (Brown, Rappert and Webster 2000), and foreclosed futures can be challenged. Whilst there are certainly categories and expectations which are applied to the Muslim community of Singapore, this does not imply that all Singaporean Muslims remain within these bounded identities. As we will explore in the next chapter, perhaps one way to resist the categories of the AI archive, and AI epistemology, is to act in contradiction of their categories.

Of course, to act outside of these categories is not always easy in Singapore. Subhas and Preeti Nair are testament to that. Therefore, subversion of the prescribed future are often subtle and opaque. Consequently, there is an ongoing debate about visibility in resistance, especially to digital technology. When computer scientist Joy Buolamwini found that her facial recognition software could not recognise her face because its training set was made up almost entirely of white faces, she began a project of creating a more inclusive dataset. The solution, she proposed was for black faces to be more visible in the data set. Something similar could be proposed for better representing Singaporean Muslims in ChatGPT's archive.

However, there is a counterargument, an argument which also surrounds more traditional archives. As Carter (2004) argues, there is silence in archives, but sometimes there is power in silence. Perhaps, Carter argues (2004), some of those not included in the archive would prefer not to be recorded, so as not to become subsumed and categorised by it. Being visible, as I return to in the next chapter, is not always desirable. The challenge, then, is how to resist foreclosed futures without becoming subsumed by them?