

Title	A LangChain-based pipeline for one-shot synthetic text generation using generative pre-trained transformers in palliative care research
Authors	Ronan, Isabel;Crowley, Patrice;Rombouts, Eva;Cornally, Nicola;Saab, Mohamad M.;Murphy, David;Tabirca, Sabin
Publication date	2025-10-15
Original Citation	Ronan, I., Crowley, P., Rombouts, E., Cornally, N., Saab, M. M., Murphy, D. and Tabirca, S. (2025) 'A LangChain-based pipeline for one-shot synthetic text generation using generative pre-trained transformers in palliative care research', Journal of Biomedical Informatics, 171, 104936 (12pp). https://doi.org/10.1016/j.jbi.2025.104936
Type of publication	Article (peer-reviewed)
Link to publisher's version	10.1016/j.jbi.2025.104936
Rights	© 2025, Elsevier Inc. For the purpose of Open Access, the author has applied a CC BY public copyright licence to any Author Accepted Manuscript version arising from this submission. - https://creativecommons.org/licenses/by/4.0/
Download date	2026-08-31 22:03:34
Item downloaded from	https://hdl.handle.net/10468/18068



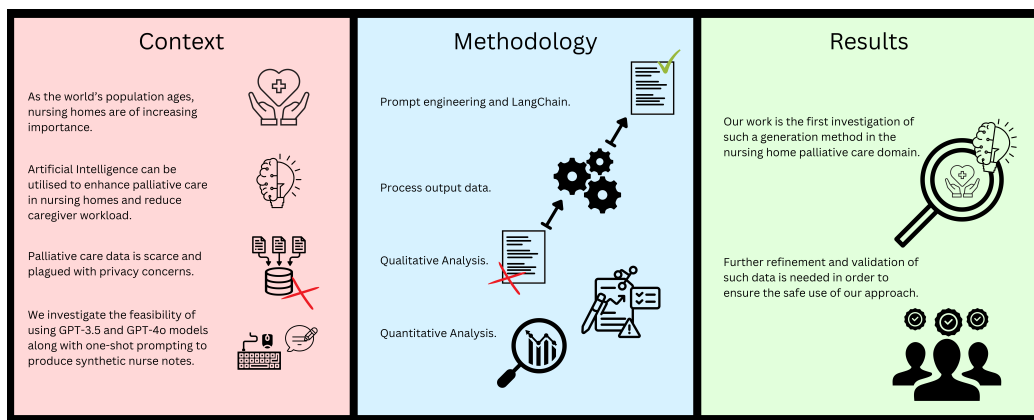
UCC

University College Cork, Ireland
Coláiste na hOllscoile Corcaigh

Graphical Abstract

A LangChain-Based Pipeline for One-Shot Synthetic Text Generation Using Generative Pre-Trained Transformers in Palliative Care Research

Isabel Ronan, Patrice Crowley, Eva Rombouts, Nicola Cornally, Mohamad M. Saab, David Murphy, Sabin Tabirca



A LangChain-Based Pipeline for One-Shot Synthetic Text Generation Using Generative Pre-Trained Transformers in Palliative Care Research

Isabel Ronan^{a,1}, Patrice Crowley^a, Eva Rombouts^b, Nicola Cornally^a,
Mohamad M. Saab^a, David Murphy^a, Sabin Tabirca^a

^a*University College Cork, College Road, Cork, T12 K8AF, Cork, Ireland*

^b*Rombouts Ouderengeneeskunde, Amsterdam, Netherlands*

Abstract

Objective: As the world’s population ages, nursing homes are of increasing importance. In order to care for a growing number of older adults, intelligent technologies are needed. Artificial Intelligence can be utilised to enhance palliative care in nursing homes. However, the data needed to train artificially intelligent agents is lacking within this sensitive domain due to privacy issues. Therefore, it is difficult for researchers to develop technological solutions. With the advent of large language models, such as ChatGPT, new text generation methods are made possible using limited data. In this pilot study, we investigate the use of large language models to generate synthetic data.

Methods: We investigate the feasibility of using GPT-3.5 and GPT-4o models along with one-shot prompting to produce synthetic nurse notes which faithfully describe nursing home residents with met or unmet palliative care needs. We used LangChain to create a repeatable pipeline which can be adapted to different use-cases. We also compare the performance of both models using a set of qualitative and quantitative evaluations to determine

¹Corresponding Author Email: i.ronan@cs.ucc.ie. Address: Western Gateway Building, University College Cork, Western Road, Cork, T12 XF62

which set of notes is more suitable for subsequent research.

Results: GPT-3.5 performed slightly better than GPT-4o in our qualitative healthcare professional analysis. Quantitative analysis revealed appropriately heterogenous results across contextual similarity, lexical overlap, sentiment, and readability scores.

Conclusion: Our work is the first investigation of such a generation method in the nursing home palliative care domain. Further refinement and validation of such data is needed in order to ensure the safe use of our approach.

Keywords: Synthetic, Palliative, Dataset, GenAI, LLM, Prompts, LangChain

1. Introduction

The world’s population is ageing rapidly. According to UN predictions, over 60 countries will have over 2 million people aged 65 or older by 2030 [1]. A significant portion of older people reside in nursing homes due to chronic diseases or other serious ailments that require constant and specialised medical care [2]. Therefore, nursing home care is growing in importance as the global population ages [3].

It is becoming increasingly difficult for limited healthcare teams to provide high quality care in a timely fashion to a growing number of older adults using traditional methods. While such care has been shown to lead to increased quality of life, the ageing population is putting pressure on resource-constrained staff within nursing homes [4]. Older adults in these facilities tend to be reliant on a small number of staff members for their care [5]. New forms of care are needed in order to ensure nursing home residents are adequately cared for.

Palliative care is an approach to care which emphasises relief from suffering and improved quality of life [6, 7, 8]. Palliative approaches have been shown to improve the quality of care received by residents in nursing homes by reducing hospital admissions, increasing clinical care access when neces-

20 sary, and ameliorating family care perceptions [9]. Palliative care can be
21 used to meet the diverse needs of older adults with appropriate impact and
22 efficiency [10]. This type of care is becoming increasingly more popular as
23 our society’s health needs change.

24 Additionally, Artificial Intelligence (AI) can be utilised to improve the
25 quality of life of nursing home residents, by providing health professionals
26 with tools to alleviate workload. AI is a broad field encompassing all systems
27 which exhibit human-like intelligence [11]. Artificially intelligent palliative
28 care solutions involving machine learning (ML) have already been successful
29 in other areas, such as in cancer patient treatment [12]. ML is a subset of
30 AI, which is often used synonymously with the broader discipline [11]. ML
31 uses algorithms coupled with large amounts of data to “learn” patterns and
32 extract meaningful information from unseen input [11]. However, researchers
33 must think carefully about the correct data to use for ML solutions in pal-
34 liative care research.

35 Palliative care is highly subjective and personalised. Subtle changes in
36 residents requiring increased care are difficult to identify. Nurse notes are
37 subjective, personalised, human-based records of resident status which are
38 collected across clinical settings and countries [13, 14]. Nurse notes may be
39 utilised to create ML models for palliative care. While ML solutions using
40 these notes could be hugely useful within nursing homes, data restrictions
41 present issues for those training algorithms in the field [15]. Furthermore, a
42 considerable quantity of nurse notes is required to train a model of sufficient
43 reliability.

44 Due to uncertainty surrounding what data to share, insufficient infrastruc-
45 ture and non-standardised multi-modal methods of data collection, ML solu-
46 tions are stagnating within the palliative nursing home domain [15]. Health-
47 care data is subject to intense privacy protection measures due to its highly
48 sensitive nature [16]. Additionally, palliative care data faces further data
49 anonymization issues due to its individualised, sparsely available nature [17].

Specifically, access to real notes from nursing homes is difficult to obtain as they are extremely personal and oftentimes unique to residents. Therefore, in order to test new approaches to machine learning within the space, a new method of obtaining big data, such as synthetic data generation (SDG), is required.

SDG is the creation of artificial data [16]. SDG circumvents many privacy and anonymization issues provided that no real patient data is used in the generation process. However, synthetic data must be plausible within the context it is generated for [16]. SDG in healthcare is fast becoming an attractive research area [18]. However, a thorough search of the literature revealed that no SDG procedures for palliative care exist, especially those for clinical note creation [17].

In this pilot study, we investigate the feasibility of using Generative Pre-Trained Transformer (GPT) models and one-shot prompting to generate synthetic notes which are representative of residents with met and unmet palliative needs in nursing homes. To generate the notes, we used LangChain, the OpenAI Application Programming Interface (API), prompt engineering techniques, and nurse-created example input. We also present the first performance comparison of GPT-3.5 and GPT-4o models for such data using both quantitative metrics and qualitative analysis by a subject matter expert. Our results show that such data needs further refinement and validation before use in the sensitive palliative care field. Additionally, comparison between synthetic and real notes is needed. The rest of this paper is divided into the following sections: literature review, methodology, results, discussion, and conclusions.

Statement of Significance

Problem or Issue	There is a lack of palliative care patient data from nursing home settings due to privacy concerns and the highly sensitive nature of end-of-life care. Lack of data limits artificial intelligence usage in these settings.
What is Already Known	GPT models can generate promising data for healthcare settings. Large language models pose problems when trained on sensitive data. Appropriate model inputs and repeatable processes are needed to generate usable synthetic data. There is a need to investigate free-text note generation specifically for palliative nursing home settings.
What this Paper Adds	This paper introduces a novel approach to palliative care nursing home data generation and evaluation using LangChain, GPT-3.5, and GPT-4o, along with both qualitative and quantitative evaluation measures.
Who would benefit from the new knowledge in this paper	Palliative care researchers, computer science researchers and healthcare professionals who have an interest in developing data-driven solutions to problems within long-term care settings.

2. Related Work

While clinical notes are crucial in many care settings, the highly individualised, subjective, free-text nature of them makes generation a unique challenge for researchers [19]. Clinical notes differ from relational data in the generation process due to the need for specialized natural language processing techniques [19]. However, as of yet, Electronic Health Records (EHR) text generation is still in its early stages [20]. Deep generative models, such as Generative Adversarial Networks (GAN), seem to be the most promising for palliative healthcare data generation at present; however, many of these

87 models tend to focus solely on tabular data creation [17].

88 Begoli et al. [21] proposed “SynthNotes”, a generator framework for high-
89 volume, high fidelity synthetic health notes. This framework operates on a
90 template-filling principle using an openly available database for information
91 extraction. Evaluation metrics for the same included BLEU and METEOR
92 scores. This system was shown to be realistic and efficient; however, it only
93 focused on the generation of mental health notes and was not extended for
94 other use-cases, such as palliative care.

95 Libbi et al. [22] proposed the use of large language models (LLMs) to
96 generate synthetic clinical notes. They used GPT-2 and real patient records
97 to produce coherent notes, comparing these notes to alternative Long Short-
98 Term Memory (LSTM) methods. A variety of evaluation metrics were used,
99 such as ROUGE-N and human evaluation, to assess the quality of their syn-
100 thetic data. Using in-text annotation, the models could produce artificial text
101 that was automatically prepared for named entity recognition (NER) tasks.
102 However, they highlight privacy concerns in their results; the model repro-
103 duced substantial portions of identifying information from the real records.
104 Additionally they describe the need for further research before EHRs can be
105 used as an anonymous alternative to real text.

106 Amin-Nejad et al. [23] used a vanilla transformer architecture to produce
107 hierarchical long passage text when trained on a large dataset. However, they
108 found the same model to be ineffective with smaller datasets. They compared
109 these results against a GPT-2 model and found that GPT-2 performed better
110 in resource-constrained environments with smaller amounts of data. They
111 also found that GPT-2 was better suited to conditional language modelling.
112 They used perplexity, BLEU and ROUGE scores to intrinsically evaluate the
113 output text. However, they used a dataset which consisted solely of notes
114 from critical care hospital admissions, thereby limiting the application of
115 their approach to palliative care settings.

116 Sufi [24] has outlined an automated pipeline for generating synthetic hos-

117 pital data using GPT-4. This work is the first which aims to use the Ope-
118 nAI API to create synthetic patient discharge messages. Results from these
119 experiments indicate that GPT models can be used to produce synthetic
120 data which is true to patient realities. However, the use of Microsoft Power
121 Automate in this method limits scalability; only 70 patient messages were
122 generated during this study. Additionally, this work does not focus on any
123 of the subtleties associated with palliative care documentation.

124 Suvalov et al. [25] investigated the use of a locally-hosted GPT-2 model
125 to generate synthetic data for further NER tasks. Their methods involved
126 the use of approximately 10 million texts from 200,00 Estonian patients to
127 generate 4100 documents for subsequent use in third-party models. However,
128 their methodology does not guarantee that input data is not leaked from the
129 synthetic data generation stages to the actual online models. Additionally,
130 palliative care is not the primary focus of this work and therefore, results are
131 not focused on capturing the subtle nuances of the field.

132 Frei and Kramer [26] have also made use of locally run GPT models to
133 produce synthetic data. They ran a GPT NeoX language model using a few-
134 shot prompting approach to generate structured German clinical text data.
135 They subsequently used their data to train a NER model for classification
136 tasks. However, this method relies on a large amount of computational power.
137 Additionally, the clinical dataset is not focused on palliative care.

138 While the aforementioned studies have used LLMs, most have not focused
139 on the creation of repeatable pipelines or frameworks to generate their infor-
140 mation. While there is little research done into the use of LLMs for synthetic
141 palliative data generation, investigations in the broader healthcare research
142 space reveal that these models pose problems when trained on sensitive data,
143 when they misunderstand their prompt input or when slight deviations in
144 prompts produce dramatically different results [27]. To solve these issues,
145 appropriate data, prompts and pipelines are needed for consistent and useful
146 synthetic data.

LangChain is an open-source framework which aims to simplify LLM application development [28]. This framework has been used by Singh et al. [29] to investigate chatbot-type tools for patients. Additionally, Kapoor and Shetty [30] proposed a chatbot focused on general health using the framework. LangChain has also been used by Ke et al. [31] to create an LLM-based preoperative aid for healthcare professionals. However, as far as we are aware, LangChain has not yet been used for synthetic data generation in healthcare.

There is a need to investigate free-text note generation specifically for palliative nursing home settings. So far, research has focused on either hospital or mental health settings. Additionally, the majority of the models discussed in the aforementioned papers use real data in the training process. The use of GPT for synthetic palliative care data generation has not been thoroughly investigated as of yet. Furthermore, one-shot prompting, the generation of synthetic data from singular synthetic notes, and LangChain have not been fully investigated. We aim to address these research gaps in our approach.

3. Methods

In this section we outline our chosen models, prompt engineering techniques, generative process, and subsequent qualitative analysis procedure. Our methodology ensures that our system is scalable and flexible, allowing for adjustable inputs and extendable outputs. We also ensure that we only used easily accessible software, services and libraries, such that minimal financial or privacy-related barriers inhibit the recreation of our research or experimental results.

3.1. Choice of Model

To demonstrate the flexibility of LangChain and to compare model outputs, we performed the data generation process with both GPT-3.5 and GPT-4o models. GPT-3.5 is a transformer-based autoregressive language model; it is capable of producing output which is difficult to distinguish

175 from human text [32, 33]. This model has a knowledge base consisting of
176 materials dated up to September, 2021 and has a 16,385-token context win-
177 dow [33]. GPT-4o is a more advanced, multi-modal transformer-based model
178 which is designed to improve on some of the shortfalls of GPT-3.5 and GPT-
179 4 [34, 33]. It is trained on data created prior to October, 2023 and has a
180 128,000-token context window [33]. At the time of writing the most recent
181 GPT-4o model released is `gpt-4o-2024-05-13`; we compare this model with
182 the `gpt-3.5-turbo-0125` GPT-3.5 model. Table 1 outlines the differences
183 between these models [33]. In this study, we apply the same prompting and
184 data “cleaning” process to the outputs of both models. Both models tem-
185 perature settings are set at 1.1; we chose this parameter due to the use and
186 validation of high temperature settings in past medical LLM research [35, 36].
187 Additionally, past research has also highlighted that statistically significant
188 results and more creative outputs are more likely at temperatures over 1.0
189 [37, 38].

Table 1: GPT-3.5 vs GPT-4o Model Comparison

	<code>gpt-3.5-turbo-0125</code>	<code>gpt-4o-2024-05-13</code>
Context Window	16,385 Tokens	128,000 tokens
Max Output Tokens	4,096 tokens	4,096 tokens
Training Data	Up to Sep 2021	Up to Oct 2023

190 3.2. Process

191 We used Google Colab and the LangChain framework to implement our
192 data generation process [39]. We used LangChain’s `PromptTemplate` compo-
193 nent to run similar prompts repeatedly [40]. In order to ensure a fluid flow
194 between system “role” and data generation prompts, we used LangChain’s
195 `SequentialChain` component [41]. Each input note was passed to a func-
196 tion which could prompt the model using both a system “role” prompt and a

197 data generation prompt. The output of the model is subsequently “cleaned”
 198 and results are saved to CSV files. Each input note is used separately, and
 199 therefore there is an individual CSV file created for each associated output
 200 batch. Figure 1 outlines our note generation pipeline.

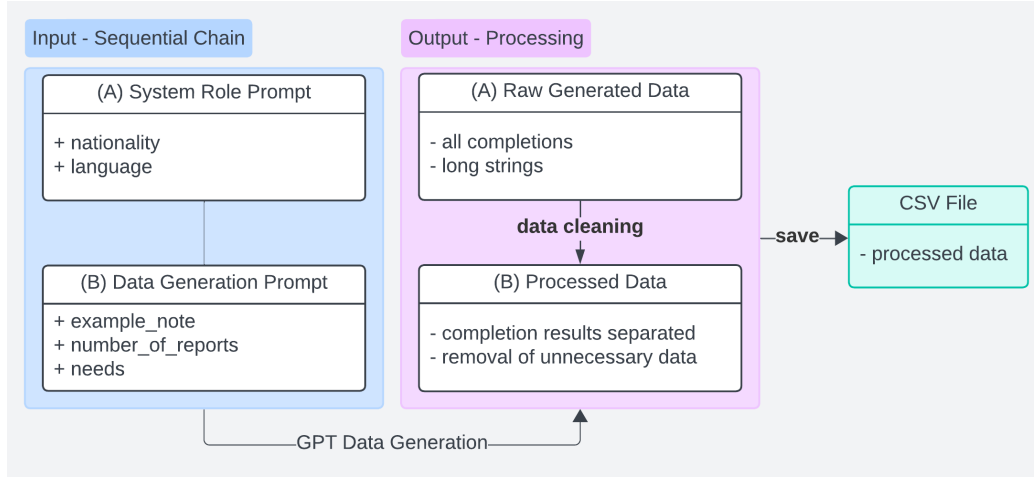


Figure 1: Note Generation Pipeline

201 3.3. Prompt Engineering

202 Prompt engineering involves the design of prompts which enable models
 203 to generate appropriate output across different domains [42]. These prompts
 204 can be optimized for specific tasks [42]. Prompt engineering is essential when
 205 interacting with LLMs to produce efficient, useful, and consistent generated
 206 output [43]. In the formulation of our prompts, we used many prompt engi-
 207 neering techniques. We also followed up-to-date guidelines provided for the
 208 medical field [43].

209 3.3.1. System Role Prompt

210 One common prompt engineering technique is to assign the model a “role”
 211 [43]. This “role” subsequently governs the model’s data generation process
 212 and allows for more structured and consistent output. We decided to make

213 the model “role” flexible using LangChain; therefore, the associated prompt
214 takes **nationality** and **language** as input variables. This flexibility allows
215 our pipeline to be implemented by someone coming from a different cultural
216 or geographical standpoint. Figure 2 describes this system “role” prompt.

```
system_role_prompt = {'message':  
    '''  
        You are a specialist in generating fictitious  
data for natural language processing projects in healthcare.  
        You speak the language of a nurse in an  
{nationality} nursing home. Namely, you speak {language}.  
    ''',  
    'inputs': ["nationality", "language"]}
```

Figure 2: System Role Prompt

217 3.3.2. Data Generation Prompt

218 In line with prompt engineering techniques recommended for healthcare,
219 we formulated the prompt input using expert knowledge [43]. Prompts were
220 crafted in terms of directness and specificity using iterative testing; this ap-
221 proach allowed for focused responses. The topics included in the prompt
222 were selected based on the content analysis of Dutch care notes using Latent
223 Dirichlet Allocation (LDA) topic modeling. Additionally, the model is told
224 not to generate overly-dramatic notes, as the majority of nursing home notes
225 contain daily activities and mundane information. Figure 3 shows this data
226 generation prompt.

227 GPT models have already proven themselves effective in resource-constrained
228 situations [23]. Therefore, we decided to use a one-shot prompting approach
229 to generate our notes, as recommended for healthcare research [43]. The
230 model is provided with an example note to replicate. Input prompt notes
231 were created by a qualified nurse. Each of these notes was designed to mimic
232 a typical progress note which would be recorded for a resident during the
233 daily documentation processes in a nursing home. The advantage of using

```

note_query_prompt = {'message':
'''
    This is an example of a nurse note for a patient in a day:
    "{example_note}"

    Other reports may include: washing, dressing, brushing
    teeth, getting ready for the day, getting ready for the
    night, showering, cleaning dental prostheses, or assistance
    after incontinence.
    Other reports could include: what the client has or has not
    eaten, what help is needed with eating (full help,
    encouragement, adapted cutlery or cup), choking, keeping
    hydration and nutrition lists.
    Other reports could include: Organised activities, getting
    visitors, browsing through a magazine, interacting with
    fellow residents. Keep in mind that these are reports from
    people in a nursing home, with severe disabilities, so
    social interaction and activities are limited. Usually it
    involves sociability, but not always.
    Other reports may include, for example: oedema, pressure
    ulcers, peeling, redness and itching of the skin. Nails that
    are too long, blemishes.
    Other reports could include, for example: care plan
    discussions, minor medical complaints, family requests,
    ordering medication.
    Reports can be, for example, about: restlessness and
    wandering at night, sleeping well, going to the toilet at
    night, phoning, lying crookedly in bed.
    Reports may include: agitation, restlessness, apathy,
    confusion; usually the confusion is subtle, but sometimes
    more intense.
    Reports may include, for example: pain, tightness of breath,
    nausea, diarrhoea, back pain, palliative care; usually the
    complaints are subtle, but sometimes more severe.
    Other reports can be about, for example: walking aids, the
    wheelchair, falls, fall incidents, transfers, lifts.
    Most reports are about everyday things, so not everything is
    a serious incident.

    Make up {number_of_reports} such reports for
    {number_of_reports} residents with {needs} palliative care
    needs. Return only the reports, with each report separated
    by "****" and nothing else. Vary the sentence structure and
    style.
''' ,
'inputs':["example_note", "number_of_reports", "needs"]}

```

Figure 3: Data Generation Prompt

234 these notes is that they are realistic, but they are not from real residents;
235 this ensures that generated data is free from privacy protection issues. Five
236 notes described residents exhibiting signs of unmet palliative care needs. An
237 additional five notes described residents not exhibiting signs of unmet pal-
238 liative care needs. We labelled these notes “Unmet” and “Met” respectively.
239 Table 2 outlines each note used in our study.

240 For the purposes of this study, we chose to execute 25 completions per
241 input note and to generate 25 notes per completion. A completion is the
242 model’s generated output; it is a prediction of what is the most likely text to
243 follow a given prompt [34]. Importantly, each completion is stateless when
244 using the OpenAI API; every fresh completion does not remember the pre-
245 vious completion, leading to more diverse output and different writing styles
246 within the notes. We decided that the model would generate 25 notes per
247 completion in order to make the overall data generation process more effi-
248 cient and diverse. However, there is no universally accepted standard in the
249 number of completions or number of data points to use when prompting a
250 model for synthetic data; every context requires a different choice[44, 45].
251 Prior work has shown that around 20-30 generations will lead to good diver-
252 sity results, provided redundancy is not too large [18]. Our high number of
253 completions was chosen with the knowledge that the models do not always
254 generate appropriate output; we chose to use a high number of completions
255 in order to select the best results from the model in later processing stages.
256 Due to the flexibility of LangChain, the number of notes and completions
257 can be adjusted to suit other use cases as future needs arise. Each of the 10
258 input notes was fed to the model individually, leading to 10 different prompt-
259 ing phases using the same prompt template. As we prompted the model to
260 generate 25 notes for each of 25 completions, the total possible number of
261 synthetic notes generated for each input note was 625. Therefore, with these
262 10 notes, our experimental setup had the potential to generate 6250 notes
263 for each model used (GPT-3.5 and GPT-4o).

Table 2: 10 Input Notes

ID	Note	Needs
1	John was assisted with a shower this morning. He had his lunch in the canteen, he ate half a bowl of soup, half a chicken salad and a full bowl of ice cream along with 2 cups of tea. Johns' sister Mary came to see him today and he was in good form.	"Met"
2	Sarah had a GP appointment this morning and the GP was happy with her mobility progress. She has been advised to continue with her physiotherapy, but her pain relief has been reduced as her pain has been well controlled. The GP hopes that she will be able to walk without crutches within the next month.	"Met"
3	Michael showered independently this morning. It was his birthday today, so his family came to visit. His wife organised a sing song in the social room. Michael was in very good spirits and sang a few songs himself.	"Met"
4	Rachel was assisted with a wash this morning, her eyedrops were given, and moisturiser applied to her legs. She went for a walk in the garden in the afternoon and attended the painting activity in the social room. Rachel has advised that she needs new pyjamas and her family have contacted about same.	"Met"
5	Darragh has been started on an End of Life care plan. There was a family conference today including Darragh, his wife Eleanor, and his children. Goals of care have been discussed and his care plan has been updated. Darragh has no complaints of pain at present, however the GP has prescribed pain relief in case there is a need. Darragh is for regular turning to avoid breakdown of his skin and his skin is fully intact at present.	"Met"
6	Billy's sister has expressed her concern that his mobility is worsening, and he is becoming increasingly dependent on others for care.	"Unmet"
7	Stephen complained of pain 8/10 in his right heel this morning. His regular analgesia was given with little effect. Visually, his heel appears normal, and his skin is intact. However, Stephen is unable to weight bear on his right leg. Stephen has no PRN analgesia prescribed, and the GP has been contacted to review same.	"Unmet"
8	Andrea experienced nausea during lunch today. She was unable to eat and was given an anti-emetic for same. She returned to bed to lie down for an hour. Afterwards, she attempted to eat some toast and vomited immediately post. She has been unable to eat or drink since. She has been commenced on subcutaneous fluids and IM Zofran has been given. Awaiting GP review.	"Unmet"
9	Louise was assisted with a wash this morning. It was discovered that she had a 5cm lesion to her left buttock that contains slough and appears infected. A swab was taken and a wound chart commenced. Inadine and Leukomed (wound dressings) in situ at present. Awaiting review from Tissue Viability Nurse.	"Unmet"
10	Michelle was assisted with her breakfast this morning by one of the healthcare assistants. The healthcare assistant expressed her concern that Michelle's swallow seems to be worsening. A referral has been made for the Speech and Language Therapist to review. For now, we have placed Michelle on a level 5 diet and level 2 fluids until she is reviewed.	"Unmet"

264 3.4. Analysis

265 We chose a variety of quantitative and qualitative measures to evaluate
 266 our synthetic output. In this section, we describe details of these measures.

267 3.4.1. Qualitative Analysis

268 We performed a qualitative analysis based on human evaluation of the
 269 notes. Three subject matter experts, reviewed and blindly labelled samples

270 of synthetic notes, which the model had generated in response to either a
271 “Met” or “Unmet” prompt. Approximately 10% of the GPT-3.5 notes and
272 10% of the GPT-4o notes were analysed. All GPT-3.5 notes were randomly
273 shuffled and 600 notes were identified for review. The same shuffling pro-
274 cess was repeated with the GPT-4o notes, and 650 notes were extracted.
275 Both sets of notes contained an approximately equal amount of “Met” and
276 “Unmet” notes. These notes were saved in Microsoft Excel files along with
277 their corresponding labels. Separate Microsoft Excel files were prepared from
278 which the labels were removed from the notes. These de-identified notes were
279 subsequently sent to the three subject matter experts for blind relabelling.
280 We also included the option to comment on each of the notes within each
281 Excel file. Each of the three evaluators was not aware of the purpose of this
282 labelling.

283 In order to effectively compare the results from our human evaluations and
284 model outputs we calculated Fleiss’ Kappa to assess the level of agreement
285 between human results [46, 47]. We subsequently used majority voting to
286 determine which human-generated labels to compare against the GPT labels.
287 Furthermore we used Cohen’s Kappa to determine the level of agreement
288 between GPT-generated notes and human-generated labels.

289 3.4.2. Quantitative Analysis

290 BertScore is a contextual similarity metric created by Zhanget al. [48]
291 which correlates better with human assessment than common text evaluation
292 metrics. BERTScore aims to account for meaning in compositionally diverse
293 text. BERTScore is calculated by computing the cosine similarity between
294 two contextualized token embeddings. BERTScore reports precision, recall
295 and F1 scores, and Zhang et al. recommend the use of the F1 score in
296 evaluations. Using the `bert-score` 0.3.13 Python library, each generated
297 note was compared against its corresponding nurse-created example note [49].
298 As BERTScore is sensitive to punctuation, we removed all punctuation from
299 our generated and comparison notes before computing BERTScore [50].

300 As outlined by Lavie and Agarwal [51], METEOR is a score used to repre-
301 sent lexical overlap. This metric is computed by matching words in an input
302 sentence with those in an output sentence. METEOR has a high correlation
303 level with human judgement. Like BERTScore, we compute METEOR met-
304 rics by comparing example input notes to generated output notes. We used
305 NLTK’s `meteor_score` module to calculate our results [52].

306 Additionally, we used the `TextBlob` Python library to calculate both sen-
307 timent and subjectivity as it has been used in other clinical literature [53].
308 `TextBlob`’s sentiment score is determined using a trained machine learning
309 model, while the subjectivity measure is calculated by evaluating the “in-
310 tensity” of each word in a sentence [54]. We compared the sentiment and
311 subjectivity scores between the nurse-generated input, GPT-3.5 and GPT-4o
312 notes using this library.

313 Due to its use in medical documentation research, we also used the Simple
314 Measure Of Gobbledygook (SMOG) index to report on the readability of
315 the generated notes [55, 56]. The `readability` Python library was used to
316 calculate the SMOG index for the GPT-3.5 and GPT-4o notes [57]. Finally,
317 we performed keyword analysis using the `wordcloud 1.9.3` Python library
318 in order to gain an overview of the most common words found in the synthetic
319 data [58].

320 We measured the time taken for both models to finish generating output
321 using a completion speed test. We executed 5 completions per input note,
322 with each completion generating 5 notes across 5 timing tests for each of
323 the two models. Each of the tests was run on a Python 3 Google Compute
324 Engine backend with 12.7GB of system RAM. Time was measured in seconds
325 using Python’s `time` library. The results of these tests were then averaged
326 to determine which model took longer to generate output.

327 We also compared our chosen one-shot prompting style, as outlined in
328 Figures 2 and 3, against a zero-shot approach. We used the BERTScore,
329 METEOR, SMOG and sentiment analysis metrics outlined above to assess

330 the quality of our prompt with and without example-related prompt text
331 given. We used 5 completions and 5 notes to compare both the original
332 prompt and the zero-shot approach. The zero-shot system-role prompt re-
333 mained the same. Figure 4 describes the data generation prompt used in the
334 zero-shot comparison scenario. We used both GPT models in the evaluation
335 of our prompt style and averaged the results across both models to create
336 comparison metrics.

```
note_query_prompt = {  
    'message':  
    '''  
    Make up {number_of_reports} nurse notes for  
    {number_of_reports} residents with {needs} palliative  
    care needs. Return only the reports, with each report  
    separated by "***" and nothing else. Vary the sentence  
    structure and style.  
    '''  
}
```

Figure 4: Zero Shot Prompt

337 4. Results

338 In this section, we report on our experimental outputs and subsequent
339 analysis. As outlined in our methodology, we used “role” formulation, di-
340 rect language, example input, and iteration prompt engineering techniques,
341 along with the OpenAI API and the LangChain framework, to generate our
342 synthetic data. All qualitative and quantitative results described below as
343 percentages are reported to two decimal places and all fractional results are
344 reported to four decimal places.

345 4.1. Model Outputs

346 There were a total of 12036 notes created from a combination of both
347 the GPT-3.5 and GPT-4o processes. Table 3 outlines detailed numerical in-
348 formation about the generated outputs, including the number of notes after
349 manual error removal. Overall, the majority of notes were preserved during
350 post-processing. While each model was capable of producing 6250 notes us-
351 ing our pipeline, this number is not fully reflected in the dataset due to minor
352 processing issues. These data processing issues involved the models produc-
353 ing output which was not properly delineated; each note was supposed to be
354 separated by “***”, but some model outputs did not include appropriate use
355 of this formatting instruction. As a result, some model outputs were saved as
356 extremely long notes; the majority of these were caught in initial processing
357 stages after the model completion and before data saving. However, the few
358 noisy entries that made it through to the final data were manually removed
359 after all model completions were saved.

360 4.2. Qualitative Analysis

361 Our human evaluation results, based on nurse-labelling of the generated
362 notes, are outlined below.

363 4.2.1. GPT-3.5 Results

364 3 of the 600 GPT-3.5 notes (0.5%) were deemed to be “problematic”.
365 “Problematic” within this context refers to notes that were thought to be too
366 difficult to classify by the subject matter experts due to file formatting issues,
367 language ambiguity, or misspellings. Problematic notes were subsequently
368 removed from this analysis, leaving 597 notes for use.

369 The Fleiss’ Kappa score obtained when all three human-raters were com-
370 pared was 0.58 for the 597 usable GPT-3.5 notes, indicating high moderate
371 agreement between all three raters. 487 of the 597 notes (81.57%) were la-
372 belled similarly by both the model and the expert. The Cohen’s Kappa
373 coefficient calculated for these notes was approximately 0.6331, indicating

Table 3: Output Note Categorization

	GPT-3.5 Before Manual Processing	GPT-3.5 After Manual Processing	GPT-4o Before Manual Processing	GPT-4o After Manual Processing
Note 1 (Met)	594	592	628	628
Note 2 (Met)	575	573	610	610
Note 3 (Met)	577	575	637	637
Note 4 (Met)	570	569	636	636
Note 5 (Met)	588	588	587	587
Note 6 (Unmet)	579	578	610	610
Note 7 (Unmet)	571	569	643	643
Note 8 (Unmet)	553	552	642	642
Note 9 (Unmet)	598	597	640	640
Note 10 (Unmet)	578	577	620	620
Total Unmet	2879	2873	3155	3155
Total Met	2904	2897	3098	3098
Total	5783	5770	6253	6253

substantial agreement between the subject matter experts and the GPT-3.5 model. The F1 score calculated based on these notes was 83.78%, indicating that the model was capable of generating valid synthetic nursing home resident notes which could be used for further analysis and research.

Figure 5 describes a confusion matrix outlining a detailed breakdown of the classification outputs. From this matrix, one can see that the experts categorised the majority of notes correctly.

4.2.2. GPT-4o Results

1 of the 650 GPT-4o notes (0.15%) presented to the experts was deemed to be “problematic”. Similarly to the GPT-3.5 analysis, this note was excluded from subsequent assessments, leaving 649 notes for further use.

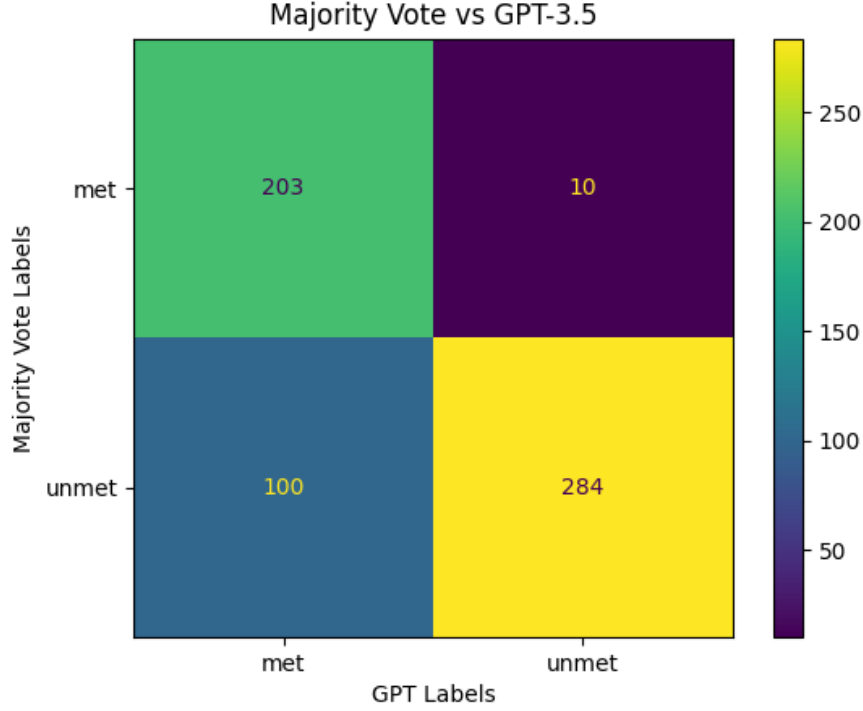


Figure 5: GPT-3.5 vs Human Experts Confusion Matrix

385 The Fleiss' Kappa score obtained when comparing our human-raters to
 386 the model was 0.299, indicating fair agreement between all three raters. 408
 387 of the 649 notes (62.87%) were labelled similarly by both the model and
 388 the expert. The Cohen's Kappa coefficient calculated for these notes was
 389 approximately 0.2674, indicating fair agreement between the subject matter
 390 experts and the GPT-4o model. The F1 score calculated based on these notes
 391 was 68.98%, indicating that the model performed poorly in comparison to
 392 the GPT-3.5 model.

393 Figure 6 describes a confusion matrix comprising the expert and GPT-4o
 394 labels. From this matrix, one can see that there is a significant disimprove-
 395 ment in the results of the GPT-4o model in comparison to the GPT-3.5
 396 model.

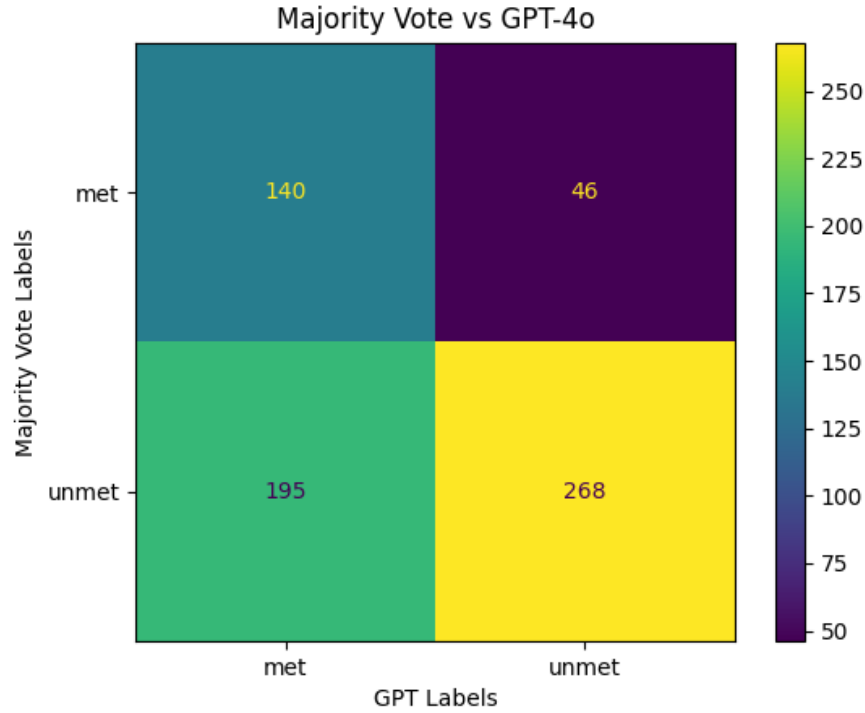


Figure 6: GPT-4o vs Human Experts Confusion Matrix

4.2.3. Overall Qualitative Results

Informal feedback from the experts indicated that the majority of notes were “quite good”. While these notes do not feature many clinical terms, the relaxed tone and informal recording of resident conditions and care seem to be in keeping with writing styles from other nursing home note analyses [59, 60]. The GPT-3.5 notes seem to be more faithful to human opinion than the GPT-4o notes when compared using the F1 and Cohen’s Kappa coefficient results of our test. Additionally, there was more inter-rater agreement in the GPT-3.5 labelling evaluation when compared against the GPT-4o evaluation. These findings indicate that, when using human-based assessments, the GPT-3.5 model is more suited to nursing home resident synthetic data creation tasks when compared against the GPT-4o model.

409 *4.3. Quantitative Analysis*

410 We performed a quantitative analysis of all notes for both the GPT-3.5
411 and GPT-4o implementations using the metrics outlined in our methodology.

412 *4.3.1. BERTScore - Contextual Similarity*

413 Table 4 outlines the BERTScore F1 results for each GPT model. While
414 somewhat contextually similar, we can see that the models produced outputs
415 that exhibit appropriate deviation from the input content. Our findings show
416 that the generated notes were sufficiently diverse enough to create appropri-
417 ately heterogeneous datasets related to the context.

Table 4: BERTScore F1 Evaluation

	GPT-3.5 (%)	GPT-4o (%)
Note 1 (Met)	51.30	51.95
Note 2 (Met)	47.65	48.2
Note 3 (Met)	48.91	50.33
Note 4 (Met)	54.16	53.53
Note 5 (Met)	55.91	48.81
Note 6 (Unmet)	49.97	45.39
Note 7 (Unmet)	51.18	52.5
Note 8 (Unmet)	48.99	50.40
Note 9 (Unmet)	47.24	49.90
Note 10 (Unmet)	47.34	48.40
Average Unmet	48.94	49.32
Average Met	51.59	50.56
Overall Average	50.27	49.94

418 *4.3.2. METEOR - Lexical Overlap*

419 METEOR scores for GPT-3.5 and GPT-4o synthetic data are outlined
420 in Table 5. These results indicate that the GPT models used a wide vocab-
421 ulary to generate varied results. This variance shows that the models did

not “overfit” to the one-shot input notes, but remained appropriately varied throughout the prompting process.

Table 5: METEOR Evaluation Reported Using Decimal Notation

	GPT-3.5	GPT-4o
Note 1 (Met)	0.1851	0.1753
Note 2 (Met)	0.1388	0.1215
Note 3 (Met)	0.1204	0.1206
Note 4 (Met)	0.2024	0.1904
Note 5 (Met)	0.2728	0.1343
Note 6 (Unmet)	0.0996	0.0712
Note 7 (Unmet)	0.1371	0.1379
Note 8 (Unmet)	0.1178	0.1192
Note 9 (Unmet)	0.0956	0.1296
Note 10 (Unmet)	0.1485	0.1524
Average Unmet	0.1197	0.1221
Average Met	0.1839	0.1484
Overall Average	0.1518	0.1352

4.3.3. *Lexicon-Based Sentiment Analysis*

Overall “Met” notes tended to be more positive than “Unmet” notes. Subjectivity results were fairly low overall, indicating that such notes are not generally subjective. Average results are reported in Table 6. From these results, we can see that both sets of model output resulted in similar sentiment and subjectivity metrics. The model results are also quite similar to those from the input nurse notes, indicating that the correct tone is reflected in these notes.

4.3.4. *Readability*

Overall, GPT-3.5 notes resulted in a SMOG index of 10.29, indicating a 10th grade reading level. GPT-4o results revealed a similar reading level with

Table 6: Sentiment and Subjectivity Reported Using Decimal Notation

	GPT-3.5	GPT-4o	Nurse Notes
Sentiment “Met” Notes	0.1798	0.1164	0.1422
Sentiment “Unmet” Notes	0.0416	0.0599	0.0733
Subjectivity “Met” Notes	0.3913	0.3922	0.3244
Subjectivity “Unmet” Notes	0.2936	0.3536	0.3650
Nurse-Model “Met” Sentiment Difference	0.0376	0.0258	-
Nurse-Model “Unmet” Sentiment Difference	0.1065	0.0430	-
Nurse-Model “Met” Subjectivity Difference	0.0669	0.0677	-
Nurse-Model “Unmet” Subjectivity Difference	0.0263	0.0272	-

435 a SMOG index of 10.69. Although the recommended reading level for medical
 436 text is 6th grade, the results from these models are in keeping with the higher
 437 grade levels of between 9 and 15 observed in freely available medical notes
 438 [61, 55]. Therefore, while simpler notes are recommended, GPT notes may
 439 be more faithful to the types of notes used in practice.

440 4.3.5. Keyword Analysis

441 The top words for each model and note category (“Met” or “Unmet”)
 442 are reported in Table 7. We removed the word “Mr” from the analysis after
 443 initial inspection as “Mr” is a shortening of “mister” and is not a word of
 444 interest per se.

445 Positive action words such as “participated” and “enjoyed” are present
 446 across both GPT-3.5 and GPT-4o “Met” notes. “Unmet” GPT-3.5 notes
 447 seem to focus on negative words such as “discomfort”, whereas GPT-4o “Un-
 448 met” notes seemed to remain fairly neutral. Figure 7 describes the overall
 449 distribution of words across all generated notes. This figure reveals a heavy
 450 emphasis on temporal words, such as “today”, “morning”, and “afternoon”.

Table 7: Top 10 Keywords for Each Model and Note Category

GPT-3.5 “Met”	GPT-3.5 “Unmet”	GPT-4o “Met”	GPT-4o “Unmet”
today	palliative care	morning	today
care plan	need	today	bed
visit	confusion	afternoon	morning
enjoyed	care team	lunch	required
afternoon	review	day	day
morning	agitation	bed	afternoon
day	restlessness	breakfast	noted
experienced	discomfort	enjoyed	night
participated	complained	participated	staff
requested	exhibited signs	seemed	wheelchair

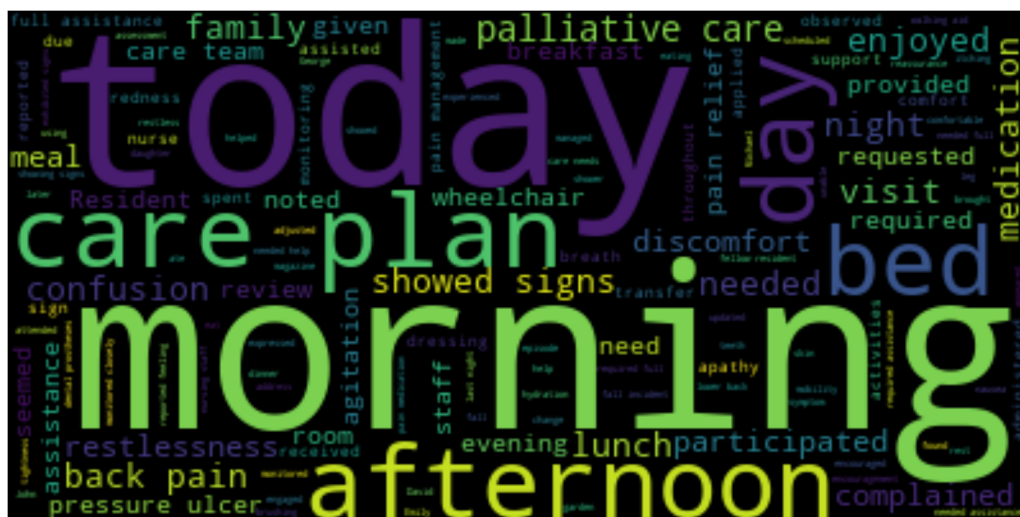


Figure 7: All Notes Word Cloud

451 4.3.6. *Timing Tests*

The results of our timing tests are outlined in Table 8. Clearly from this table, GPT-3.5 is a lot faster at completing the note generation task than

454 GPT-4o, indicating that the GPT-3.5 model is more efficient than its more
 455 recent counterpart.

Table 8: Model Completion Time

Run Number	GPT-3.5 (seconds)	GPT-4o (seconds)
Run 1	223.39	1212.10
Run 2	239.47	1085.89
Run 3	287.74	997.09
Run 4	217.71	1118.24
Run 5	211.34	1117.94
Average Time	235.93	1106.25

456 4.3.7. Prompt Comparisons

457 The results of our prompt style assessments are outlined in Table 9. From
 458 this table, it is clear that our main one-shot approach outperforms the zero-
 459 shot approach across metrics. The SMOG index results indicate that the
 460 one-shot approach results in notes of a more manageable reading level than
 461 the zero-shot approach. The BERTScore results indicate that the one-shot
 462 approach results in notes that are more similar to our nurse-generated input
 463 notes. The METEOR score also conveys the slightly better performance of
 464 the one-shot approach. There are also greater differences in subjectivity and
 465 sentiment across both met and unmet notes for zero-shot results in com-
 466 parison to one-shot results, indicating that the one-shot approach results in
 467 notes that are more similar to the input notes. While the majority of results
 468 demonstrate only small improvements across metrics, these findings indicate
 469 that the one-shot guidance results in notes that are more appropriate to the
 470 context.

471 Additionally, the zero-shot approach resulted in many notes being incor-
 472 rectly generated without proper delineation. The zero-shot prompts resulted

473 in a total of 119 notes. The one-shot prompts resulted in a total of 480 notes.
 474 Given that the total number of notes expected given our 5-note 5-completion
 475 test parameters was 500, these results indicate that the one-shot approach
 476 generated more appropriate output in this scenario.

Table 9: Prompt Style Comparison

Metric	Average Zero-Shot	Average One-Shot
SMOG	12.27	11.67
BERTScore	0.44	0.54
METEOR	0.07	0.16
Sentiment "Met" Difference	0.11	0.05
Sentiment "Unmet" Difference	0.04	0.03
Subjectivity "Met" Difference	0.1	0.09
Subjectivity "Unmet" Difference	0.13	0.025

477 5. Discussion

478 This paper presents a template LangChain and subsequent validation
 479 method for generating synthetic data for nursing home palliative care re-
 480 search. This template can be easily modified to include different topics as
 481 needed. As a result of our implementation, we were able to generate over
 482 12,000 notes using just 10 nurse-provided examples. We investigated the use
 483 of such a template to produce notes which are reflective of human-generated
 484 notes. These notes were evaluated using a variety of qualitative and quan-
 485 titative measures. We assessed the feasibility of our approach using these
 486 measures. We also compared GPT-3.5 and GPT-4o models using these as-
 487 sessments.

488 Both the GPT-3.5 and GPT-4o models exhibit similar results when com-
 489 pared using our quantitative assessments. Keyword analysis revealed that
 490 both models use temporal words. As described in Section 4.3.1, we used

491 BERTScore to compute contextual similarity; we found both models capa-
492 ble of generating appropriately contextually-diverse data. METEOR was
493 used to examine lexical overlap, as outlined in Section 4.3.2; there is suit-
494 able variance in the generated data. Section 4.3.4 outlines the result of
495 our SMOG readability assessments; our notes were of similar reading level
496 to other human-written nurse notes. Additionally, lexicon-based sentiment
497 analysis, which is described in Section 4.3.3, revealed that our generated
498 content reflects similar sentiment to our nurse-generated input notes.

499 While the models convey similar results when observed using quantita-
500 tive assessments, our qualitative evaluation reveals that GPT-3.5 notes are
501 more in-keeping with expert opinion than GPT-4o notes. We conducted a
502 blind relabelling test to find that the GPT-3.5 output is more reflective of
503 human generated notes than the GPT-4o output. Given the importance of
504 the qualitative nature of nurse notes, we posit that GPT-3.5 is more suit-
505 able for the note generation task. Furthermore, upon assessing the amount
506 of time taken to complete the note generation task, GPT-3.5 outperformed
507 GPT-4o by a significant margin, demonstrating the efficiency of the former
508 model in comparison to the latter.

509 Furthermore, GPT-4o’s ability to handle lengthier text with increased
510 complexity may be disadvantageous in simple and concise note generation
511 contexts. Through these results, we show that more recent models do not
512 necessarily outperform older models across all metrics or evaluations. This
513 finding implies that, when generating synthetic data, multiple models should
514 be considered, not just the most recent versions.

515 This study is not without limitations. All of the nurse notes used as
516 example input during testing were generated by one nurse. The nurse expert
517 has been trained and works in the Irish healthcare system only and has no
518 experience of other healthcare systems. These generated notes may not be
519 relevant outside an Irish context. Future experiments could include notes
520 from a wider variety of health professionals to further confirm the feasibility

521 of our approach.

522 While we attempted to evaluate our notes using both qualitative and
523 quantitative analysis, there exists no standard evaluation for synthetically
524 generated text [62]. While human evaluation is considered the gold stan-
525 dard, there are no specific criteria outlined for such assessments [62]. This
526 lack of standardisation extends to quantitative evaluation scenarios also. The
527 context of this work makes it difficult to compare our findings to other syn-
528 thetic data studies. Additionally, in life-and-death settings, higher standards
529 are required, rather than relying on arbitrary assessment techniques.

530 Our human evaluation is carried out on a 10% sample of randomly shuffled
531 generated notes due to time and resource constraints. While such assessment
532 is useful in determining a general overview of data quality, this approach
533 does not guarantee that our results are representative of the entire data
534 domain. Additionally, one of our evaluators generated the input notes for the
535 experiment. While this evaluator was blind to the purpose of the evaluation
536 and was not involved in the design of the experimental procedure, the use
537 of the same nurse in both generation and evaluation processes could have
538 potentially introduced bias into our study.

539 Furthermore, in an effort to discretize the palliative care referral land-
540 scape, we classify notes into either “Met” or “Unmet” data points. In pal-
541 liative care practice, it is more common to perceive such data as indicative
542 or not indicative of signs of unmet palliative care needs. Therefore, our
543 human evaluation may have resulted in an assessment skewed towards the
544 identification of unmet palliative care needs. Furthermore, discretizing such
545 a complex field may potentially hinder future use of such data. A more
546 structured approach to labelling notes with pre-defined criteria may improve
547 results.

548 While our synthetic data usage aims to ensure that sensitive information
549 from real patients is not revealed, we cannot guarantee that the foundation
550 models being used do not leak training data. LLMs are trained on a wide

551 variety of information from diverse sources; we do not currently have access
552 to the exact data that these models were trained on. Additionally, there
553 is no definitive way to check for the leakage of sensitive information from
554 these models due to the lack of knowledge surrounding the initial foundation
555 model training data. Future work could focus on using open-source models
556 with more clarity surrounding leakage risks to ensure safer synthetic dataset
557 generation processes.

558 Our results demonstrate the feasibility of using LangChain, the OpenAI
559 API, and one-shot prompting to generate synthetic data for future research.
560 However, a further blind study is needed involving a panel of experts in order
561 to validate this data for future use. If such synthetic data is valid, researchers
562 can use it in preliminary investigation stages before trying to access real-
563 world notes. These notes could be used to test potential algorithms before
564 real resident data is accessed and used by researchers. Further work could
565 compare the performance of machine learning models trained on synthetically
566 generated data in comparison to the same models trained on real resident
567 data when used on unseen data. This future assessment could inform future
568 research on the functionality of synthetic data.

569 6. Conclusion

570 This pilot study outlines a novel synthetic data generation pipeline to
571 investigate the feasibility of using GPT models to generate nursing home
572 resident data for palliative care research. We also compare results between
573 these models. Our findings indicate that such approaches to artificial data
574 creation can result in synthetic notes which may be reflective of human-
575 generated data. We also find that GPT-3.5 may be more suited to such data
576 generation tasks when compared against GPT-4o using human evaluation.
577 In this paper, we discuss our methodology, analysis, and future directions.
578 We also outline potential limitations associated with our work. Such data
579 has the potential to be used by future researchers to accelerate palliative care

580 research and to preserve patient privacy.

581 **Appendix A. Supplementary Data**

582 The code used and datasets generated during this research can be found
583 on GitHub at <https://github.com/isabel-ronan/SyntheticNotes>.

584 **CRedit Authorship Contribution Statement**

585 **Isabel Ronan:** Writing – review & editing, Writing – original draft,
586 Software, Methodology, Investigation, Conceptualization, Validation, For-
587 mal Analysis, Data Curation. **Patrice Crowley:** Writing – review & edit-
588 ing, Investigation, Conceptualization, Validation. **Eva Rombouts:** Writ-
589 ing – review & editing, Conceptualization, Validation. **Nicola Cornally:**
590 Writing – review & editing, Supervision, Conceptualization. **Mohamad M**
591 **Saab:** Writing – review & editing, Supervision, Conceptualization, Vali-
592 dation. **David Murphy:** Writing – review & editing, Supervision, Con-
593 ceptualization. **Sabin Tabirca:** Writing – review & editing, Supervision,
594 Conceptualization.

595 **Declaration of Competing Interest**

596 The authors declare that they have no known competing financial inter-
597 ests or personal relationships that could have appeared to influence the work
598 reported in this paper.

599 **Funding Sources**

600 This publication has emanated from research conducted with the finan-
601 cial support of Taighde Éireann – Research Ireland under Grant number
602 18/CRT/6222. For the purpose of Open Access, the author has applied a
603 CC BY public copyright licence to any Author Accepted Manuscript version
604 arising from this submission.

605 Patrice Crowley is in receipt of a PhD scholarship from the Catherine McAuley
606 School of Nursing and Midwifery, University College Cork, Cork, Ireland.

References

- [1] J. L. Powell, The power of global aging 35 (1) 1–14. doi:10.1007/s12126-010-9051-6.
URL <https://doi.org/10.1007/s12126-010-9051-6>
- [2] A. J. De Veer, A. Kerkstra, Feeling at home in nursing homes 35 (3) 427–434, eprint: <https://onlinelibrary.wiley.com/doi/pdf/10.1046/j.1365-2648.2001.01858.x>. doi:10.1046/j.1365-2648.2001.01858.x.
URL <https://onlinelibrary.wiley.com/doi/abs/10.1046/j.1365-2648.2001.01858.x>
- [3] L. Schenk, R. Meyer, A. Behr, A. Kuhlmei, M. Holzhausen, Quality of life in nursing homes: results of a qualitative resident survey 22 (10) 2929–2938. doi:10.1007/s11136-013-0400-2.
URL <https://doi.org/10.1007/s11136-013-0400-2>
- [4] I. Hjaltadóttir, M. Gústafsdóttir, Quality of life in nursing homes: perception of physically frail elderly residents 21 (1) 48–55. doi:10.1111/j.1471-6712.2007.00434.x.
URL <https://pubmed.ncbi.nlm.nih.gov/17428214/>
- [5] S. Hall, S. Longhurst, I. Higginson, Living and dying with dignity: a qualitative study of the views of older people in nursing homes 38 (4) 411–416. doi:10.1093/ageing/afp069.
URL <https://doi.org/10.1093/ageing/afp069>
- [6] Topics in palliative care.
URL <http://ndl.ethernet.edu.et/bitstream/123456789/45753/1/3228.pdf>

- [7] B. Wallerstedt, E. Benzein, K. Schildmeijer, A. Sandgren, What is palliative care? perceptions of healthcare professionals 33 (1) 77–84. doi:10.1111/scs.12603.
- [8] W. Van Mechelen, B. Aertgeerts, K. De Ceulaer, B. Thoonsen, M. Vermandere, F. Warmenhoven, E. Van Rijswijk, J. De Lepeleire, Defining the palliative care patient: A systematic review 27 (3) 197–208, publisher: SAGE Publications Ltd STM. doi:10.1177/0269216311435268. URL <https://doi.org/10.1177/0269216311435268>
- [9] N. M. Cimino, M. L. McPherson, Evaluating the impact of palliative or hospice care provided in nursing homes 40 (10) 10–14. doi:10.3928/00989134-20140909-01.
- [10] R. Schweighart, J. L. O’Sullivan, M. Klemmt, A. Teti, S. Neuderth, Wishes and needs of nursing home residents: A scoping review 10 (5) 854. doi:10.3390/healthcare10050854. URL <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC9140474/>
- [11] J. M. Helm, A. M. Swiergosz, H. S. Haeberle, J. M. Karnuta, J. L. Schaffer, V. E. Krebs, A. I. Spitzer, P. N. Ramkumar, Machine learning and artificial intelligence: Definitions, applications, and future directions 13 (1) 69–76. doi:10.1007/s12178-020-09600-8. URL <https://doi.org/10.1007/s12178-020-09600-8>
- [12] V. Reddy, A. Nafees, S. Raman, Recent advances in artificial intelligence applications for supportive and palliative care in cancer patients 17 (2) 125–134. doi:10.1097/SPC.0000000000000645.
- [13] G. Hansebo, M. Kihlgren, G. Ljunggren, Review of nursing documentation in nursing home wards — changes after intervention for individualized care 29 (6) 1462–1473, eprint: <https://onlinelibrary.wiley.com/doi/pdf/10.1046/j.1365-2648.1999.01034.x>. doi:10.1046/j.1365-2648.1999.01034.x.

URL <https://onlinelibrary.wiley.com/doi/abs/10.1046/j.1365-2648.1999.01034.x>

- [14] P. Voutilainen, A. Isola, S. Muurinen, Nursing documentation in nursing homes – state-of-the-art and implications for quality improvement 18 (1) 72–81, eprint: <https://onlinelibrary.wiley.com/doi/pdf/10.1111/j.1471-6712.2004.00265.x>. doi:10.1111/j.1471-6712.2004.00265.x.
URL <https://onlinelibrary.wiley.com/doi/abs/10.1111/j.1471-6712.2004.00265.x>
- [15] A. Hendriks, C. Hacking, H. Verbeek, S. Aarts, Data science techniques to gain novel insights into quality of care: a scoping review of long-term care for older adults 2 (2) 67–85, number: 2 Publisher: Open Exploration. doi:10.37349/edht.2024.00012.
URL <https://www.explorationpub.com/Journals/edht/Article/101112>
- [16] A. Figueira, B. Vaz, Survey on synthetic data generation, evaluation methods and GANs 10 (15) 2733, number: 15 Publisher: Multidisciplinary Digital Publishing Institute. doi:10.3390/math10152733.
URL <https://www.mdpi.com/2227-7390/10/15/2733>
- [17] W. Hahn, K. Schütte, K. Schultz, O. Wolkenhauer, M. Sedlmayr, U. Schuler, M. Eichler, S. Bej, M. Wolfien, Contribution of synthetic data generation towards an improved patient stratification in palliative care 12 (8) 1278, number: 8 Publisher: Multidisciplinary Digital Publishing Institute. doi:10.3390/jpm12081278.
URL <https://www.mdpi.com/2075-4426/12/8/1278>
- [18] R. J. Chen, M. Y. Lu, T. Y. Chen, D. F. K. Williamson, F. Mahmood, Synthetic data in machine learning for medicine and healthcare 5 (6) 493–497, publisher: Nature Publishing Group. doi:10.1038/s41551-021-

00751-8.

URL <https://www.nature.com/articles/s41551-021-00751-8>

- [19] H. Murtaza, M. Ahmed, N. F. Khan, G. Murtaza, S. Zafar, A. Bano, Synthetic data generation: State of the art in health care domain 48 100546. doi:10.1016/j.cosrev.2023.100546.
URL <https://www.sciencedirect.com/science/article/pii/S1574013723000138>
- [20] D. Rankin, M. Black, R. Bond, J. Wallace, M. Mulvenna, G. Epelde, Reliability of supervised machine learning using synthetic data in health care: Model to preserve privacy for data sharing 8 (7) e18910, company: JMIR Medical Informatics Distributor: JMIR Medical Informatics Institution: JMIR Medical Informatics Label: JMIR Medical Informatics Publisher: JMIR Publications Inc., Toronto, Canada. doi:10.2196/18910.
URL <https://medinform.jmir.org/2020/7/e18910>
- [21] E. Begoli, K. Brown, S. Srinivas, S. Tamang, SynthNotes: A generator framework for high-volume, high-fidelity synthetic mental health notes, in: 2018 IEEE International Conference on Big Data (Big Data), pp. 951–958. doi:10.1109/BigData.2018.8621981.
URL <https://ieeexplore.ieee.org/document/8621981>
- [22] C. A. Libbi, J. Trienes, D. Trieschnigg, C. Seifert, Generating synthetic training data for supervised de-identification of electronic health records 13 (5) 136, number: 5 Publisher: Multidisciplinary Digital Publishing Institute. doi:10.3390/fi13050136.
URL <https://www.mdpi.com/1999-5903/13/5/136>
- [23] A. Amin-Nejad, J. Ive, S. Velupillai, Exploring transformer text generation for medical dataset augmentation, in: N. Calzolari, F. Béchet, P. Blache, K. Choukri, C. Cieri, T. Declerck, S. Goggi, H. Isahara,

- B. Maegaard, J. Mariani, H. Mazo, A. Moreno, J. Odijk, S. Piperidis (Eds.), Proceedings of the Twelfth Language Resources and Evaluation Conference, European Language Resources Association, pp. 4699–4708. URL <https://aclanthology.org/2020.lrec-1.578>
- [24] F. Sufi, Addressing data scarcity in the medical domain: A GPT-based approach for synthetic data generation and feature extraction 15 (5) 264, number: 5 Publisher: Multidisciplinary Digital Publishing Institute. doi:10.3390/info15050264. URL <https://www.mdpi.com/2078-2489/15/5/264>
- [25] H. Šuvalov, M. Lepson, V. Kukk, M. Malk, N. Ilves, H.-A. Kuulmets, R. Kolde, Using synthetic health care data to leverage large language models for named entity recognition: Development and validation study 27 (1) e66279, company: Journal of Medical Internet Research Distributor: Journal of Medical Internet Research Institution: Journal of Medical Internet Research Label: Journal of Medical Internet Research Publisher: JMIR Publications Inc., Toronto, Canada. doi:10.2196/66279. URL <https://www.jmir.org/2025/1/e66279>
- [26] J. Frei, F. Kramer, Annotated dataset creation through general purpose language models for non-english medical NLP. arXiv:2208.14493 [cs], doi:10.48550/arXiv.2208.14493. URL <http://arxiv.org/abs/2208.14493>
- [27] K. He, R. Mao, Q. Lin, Y. Ruan, X. Lan, M. Feng, E. Cambria, A survey of large language models for healthcare: from data, technology, and applications to accountability and ethics. arXiv:2310.05694 [cs], doi:10.48550/arXiv.2310.05694. URL <http://arxiv.org/abs/2310.05694>
- [28] H. Chase, LangChain, original-date: 2022-10-17T02:58:36Z. URL <https://github.com/langchain-ai/langchain>

- [29] A. Singh, A. Ehtesham, S. Mahmud, J.-H. Kim, Revolutionizing mental health care through LangChain: A journey with a large language model, in: 2024 IEEE 14th Annual Computing and Communication Workshop and Conference (CCWC), pp. 0073–0078. doi:10.1109/CCWC60891.2024.10427865.
URL <https://ieeexplore.ieee.org/abstract/document/10427865>
- [30] A. Kapoor, S. D. Shetty, Enhancing healthcare information accessibility through a generative medical chatbot, in: 2024 International Conference on Emerging Technologies in Computer Science for Interdisciplinary Applications (ICETCS), IEEE, pp. 1–4. doi:10.1109/ICETCS61022.2024.10543529.
URL <https://ieeexplore.ieee.org/document/10543529/>
- [31] Y. H. Ke, L. Jin, K. Elangovan, H. R. Abdullah, N. Liu, A. T. H. Sia, C. R. Soh, J. Y. M. Tung, J. C. L. Ong, D. S. W. Ting, Development and testing of retrieval augmented generation in large language models - a case study report.
- [32] T. B. Brown, B. Mann, N. Ryder, M. Subbiah, J. Kaplan, P. Dhariwal, A. Neelakantan, P. Shyam, G. Sastry, A. Askell, S. Agarwal, A. Herbert-Voss, G. Krueger, T. Henighan, R. Child, A. Ramesh, D. M. Ziegler, J. Wu, C. Winter, C. Hesse, M. Chen, E. Sigler, M. Litwin, S. Gray, B. Chess, J. Clark, C. Berner, S. McCandlish, A. Radford, I. Sutskever, D. Amodei, Language models are few-shot learners. arXiv:2005.14165 [cs], doi:10.48550/arXiv.2005.14165.
URL <http://arxiv.org/abs/2005.14165>
- [33] OpenAI platform.
URL <https://platform.openai.com>
- [34] OpenAI, J. Achiam, S. Adler, S. Agarwal, L. Ahmad, I. Akkaya, F. L. Aleman, D. Almeida, J. Altenschmidt, S. Altman, S. Anadkat, R. Avila,

I. Babuschkin, S. Balaji, V. Balcom, P. Baltescu, H. Bao, M. Bavarian,
 J. Belgum, I. Bello, J. Berdine, G. Bernadett-Shapiro, C. Berner, L. Bog-
 donoff, O. Boiko, M. Boyd, A.-L. Brakman, G. Brockman, T. Brooks,
 M. Brundage, K. Button, T. Cai, R. Campbell, A. Cann, B. Carey,
 C. Carlson, R. Carmichael, B. Chan, C. Chang, F. Chantzis, D. Chen,
 S. Chen, R. Chen, J. Chen, M. Chen, B. Chess, C. Cho, C. Chu,
 H. W. Chung, D. Cummings, J. Currier, Y. Dai, C. Decareaux, T. De-
 gry, N. Deutsch, D. Deville, A. Dhar, D. Dohan, S. Dowling, S. Dun-
 ning, A. Ecoffet, A. Eleti, T. Eloundou, D. Farhi, L. Fedus, N. Felix,
 S. P. Fishman, J. Forte, I. Fulford, L. Gao, E. Georges, C. Gibson,
 V. Goel, T. Gogineni, G. Goh, R. Gontijo-Lopes, J. Gordon, M. Graf-
 stein, S. Gray, R. Greene, J. Gross, S. S. Gu, Y. Guo, C. Hallacy, J. Han,
 J. Harris, Y. He, M. Heaton, J. Heidecke, C. Hesse, A. Hickey, W. Hickey,
 P. Hoeschele, B. Houghton, K. Hsu, S. Hu, X. Hu, J. Huizinga, S. Jain,
 S. Jain, J. Jang, A. Jiang, R. Jiang, H. Jin, D. Jin, S. Jomoto, B. Jonn,
 H. Jun, T. Kaftan, Kaiser, A. Kamali, I. Kanitscheider, N. S. Keskar,
 T. Khan, L. Kilpatrick, J. W. Kim, C. Kim, Y. Kim, J. H. Kirchner,
 J. Kiros, M. Knight, D. Kokotajlo, Kondraciuk, A. Kondrich,
 A. Konstantinidis, K. Kopic, G. Krueger, V. Kuo, M. Lampe, I. Lan,
 T. Lee, J. Leike, J. Leung, D. Levy, C. M. Li, R. Lim, M. Lin, S. Lin,
 M. Litwin, T. Lopez, R. Lowe, P. Lue, A. Makanju, K. Malfacini,
 S. Manning, T. Markov, Y. Markovski, B. Martin, K. Mayer, A. Mayne,
 B. McGrew, S. M. McKinney, C. McLeavey, P. McMillan, J. McNeil,
 D. Medina, A. Mehta, J. Menick, L. Metz, A. Mishchenko, P. Mishkin,
 V. Monaco, E. Morikawa, D. Mossing, T. Mu, M. Murati, O. Murk,
 D. Mély, A. Nair, R. Nakano, R. Nayak, A. Neelakantan, R. Ngo,
 H. Noh, L. Ouyang, C. O’Keefe, J. Pachocki, A. Paino, J. Palermo,
 A. Pantuliano, G. Parascandolo, J. Parish, E. Parparita, A. Passos,
 M. Pavlov, A. Peng, A. Perelman, F. d. A. B. Peres, M. Petrov, H. P.
 d. O. Pinto, Michael, Pokorny, M. Pokrass, V. H. Pong, T. Powell,

- A. Power, B. Power, E. Proehl, R. Puri, A. Radford, J. Rae, A. Ramesh, C. Raymond, F. Real, K. Rimbach, C. Ross, B. Rotsted, H. Roussez, N. Ryder, M. Saltarelli, T. Sanders, S. Santurkar, G. Sastry, H. Schmidt, D. Schnurr, J. Schulman, D. Selsam, K. Sheppard, T. Sherbakov, J. Shieh, S. Shoker, P. Shyam, S. Sidor, E. Sigler, M. Simens, J. Sitkin, K. Slama, I. Sohl, B. Sokolowsky, Y. Song, N. Staudacher, F. P. Such, N. Summers, I. Sutskever, J. Tang, N. Tezak, M. B. Thompson, P. Tillet, A. Tootoonchian, E. Tseng, P. Tuggle, N. Turley, J. Tworek, J. F. C. Uribe, A. Vallone, A. Vijayvergiya, C. Voss, C. Wainwright, J. J. Wang, A. Wang, B. Wang, J. Ward, J. Wei, C. J. Weinmann, A. Welihinda, P. Welinder, J. Weng, L. Weng, M. Wiethoff, D. Willner, C. Winter, S. Wolrich, H. Wong, L. Workman, S. Wu, J. Wu, M. Wu, K. Xiao, T. Xu, S. Yoo, K. Yu, Q. Yuan, W. Zaremba, R. Zellers, C. Zhang, M. Zhang, S. Zhao, T. Zheng, J. Zhuang, W. Zhuk, B. Zoph, GPT-4 technical report. arXiv:2303.08774 [cs], doi:10.48550/arXiv.2303.08774. URL <http://arxiv.org/abs/2303.08774>
- [35] M. Tornqvist, J.-D. Zucker, T. Fauvel, N. Lambert, M. Berthelot, A. Movschin, A text-to-tabular approach to generate synthetic patient data using LLMs. arXiv:2412.05153 [cs], doi:10.48550/arXiv.2412.05153. URL <http://arxiv.org/abs/2412.05153>
- [36] D. Smolyak, A. Welivita, M. V. Bjarnadóttir, R. Agarwal, Improving equity in health modeling with GPT4-turbo generated synthetic data: A comparative study. arXiv:2412.16335 [cs], doi:10.48550/arXiv.2412.16335. URL <http://arxiv.org/abs/2412.16335>
- [37] L. Li, L. Sleem, N. Gentile, G. Nichil, R. State, Exploring the impact of temperature on large language models: hot or cold? arXiv:2506.07295 [cs], doi:10.48550/arXiv.2506.07295. URL <http://arxiv.org/abs/2506.07295>

- [38] M. Renze, E. Guven, The effect of sampling temperature on problem solving in large language models, pp. 7346–7356. arXiv:2402.05201 [cs], doi:10.18653/v1/2024.findings-emnlp.432.
URL <http://arxiv.org/abs/2402.05201>
- [39] colab.google.
URL <https://colab.google/>
- [40] langchain_core.prompts.prompt.PromptTemplate — LangChain 0.2.7.
URL https://api.python.langchain.com/en/latest/prompts/langchain_core.prompts.prompt.PromptTemplate.html
- [41] langchain.chains.sequential.SequentialChain — LangChain 0.2.7.
URL <https://api.python.langchain.com/en/latest/chains/langchain.chains.sequential.SequentialChain.html>
- [42] J. Wang, E. Shi, S. Yu, Z. Wu, C. Ma, H. Dai, Q. Yang, Y. Kang, J. Wu, H. Hu, C. Yue, H. Zhang, Y. Liu, Y. Pan, Z. Liu, L. Sun, X. Li, B. Ge, X. Jiang, D. Zhu, Y. Yuan, D. Shen, T. Liu, S. Zhang, Prompt engineering for healthcare: Methodologies and applications. arXiv:2304.14670 [cs], doi:10.48550/arXiv.2304.14670.
URL <http://arxiv.org/abs/2304.14670>
- [43] B. Meskó, Prompt engineering as an important emerging skill for medical professionals: Tutorial 25 (1) e50638, company: Journal of Medical Internet Research Distributor: Journal of Medical Internet Research Institution: Journal of Medical Internet Research Label: Journal of Medical Internet Research Publisher: JMIR Publications Inc., Toronto, Canada. doi:10.2196/50638.
URL <https://www.jmir.org/2023/1/e50638>
- [44] H. Chen, A. Waheed, X. Li, Y. Wang, J. Wang, B. Raj, M. I. Abdin, On the diversity of synthetic data and its impact on training large language

- models. arXiv:2410.15226 [cs], doi:10.48550/arXiv.2410.15226.
URL <http://arxiv.org/abs/2410.15226>
- [45] Y. Zhu, H. Zhong, Q. Lin, H. Wei, X. Sun, Z. Yu, M. Liu, Z. Zheng, L. Chen, What matters in LLM-generated data: Diversity and its effect on model fine-tuning. arXiv:2506.19262 [cs], doi:10.48550/arXiv.2506.19262.
URL <http://arxiv.org/abs/2506.19262>
- [46] T. R. Nichols, P. M. Wisner, G. Cripe, L. Gulabchand, Putting the kappa statistic to use 13 (3) 57–61, eprint: <https://onlinelibrary.wiley.com/doi/pdf/10.1002/qaj.481>. doi:10.1002/qaj.481.
URL <https://onlinelibrary.wiley.com/doi/abs/10.1002/qaj.481>
- [47] J. R. Landis, G. G. Koch, The measurement of observer agreement for categorical data 33 (1) 159–174, publisher: International Biometric Society. doi:10.2307/2529310.
URL <https://www.jstor.org/stable/2529310>
- [48] T. Zhang, V. Kishore, F. Wu, K. Q. Weinberger, Y. Artzi, BERTScore: Evaluating text generation with BERT. arXiv:1904.09675 [cs], doi:10.48550/arXiv.1904.09675.
URL <http://arxiv.org/abs/1904.09675>
- [49] Artzi, Yoav, T. Zhang, V. Kishore, F. Wu, K. Q. Weinberger, bert-score: PyTorch implementation of BERT score.
URL https://github.com/Tiiiger/bert_score
- [50] P. Jwalapuram, Pulling out all the full stops: Punctuation sensitivity in neural machine translation and evaluation, in: A. Rogers, J. Boyd-Graber, N. Okazaki (Eds.), Findings of the Association for Computa-

- tional Linguistics: ACL 2023, Association for Computational Linguistics, pp. 6116–6130. doi:10.18653/v1/2023.findings-acl.381.
URL <https://aclanthology.org/2023.findings-acl.381>
- [51] A. Lavie, A. Agarwal, Meteor: an automatic metric for MT evaluation with high levels of correlation with human judgments, in: Proceedings of the Second Workshop on Statistical Machine Translation - StatMT '07, Association for Computational Linguistics, pp. 228–231. doi:10.3115/1626355.1626389.
URL <http://portal.acm.org/citation.cfm?doid=1626355.1626389>
- [52] NLTK :: nltk.translate.meteor_score module.
URL https://www.nltk.org/api/nltk.translate.meteor_score.html
- [53] K. Denecke, D. Reichenpfader, Sentiment analysis of clinical narratives: A scoping review 140 104336. doi:10.1016/j.jbi.2023.104336.
URL <https://www.sciencedirect.com/science/article/pii/S1532046423000576>
- [54] S. Loria, sloria/TextBlob, original-date: 2013-06-30T18:29:18Z.
URL <https://github.com/slوريا/TextBlob>
- [55] A. S. Hedman, Using the SMOG formula to revise a health-related document 39 (1) 61–64. doi:10.1080/19325037.2008.10599016.
- [56] P. Fitzsimmons, B. Michael, J. Hulley, G. Scott, A readability assessment of online parkinson’s disease information 40 (4) 292–296. doi:10.4997/JRCPE.2010.401.
URL http://www.rcpe.ac.uk/journal/issue/journal_40_4/fitzsimmons.pdf

- [57] A. v. Cranenburgh, readability: Measure the readability of a given text using surface characteristics.
URL <https://github.com/andreascv/readability/>
- [58] A. Mueller, wordcloud: A little word cloud generator.
URL https://github.com/amueller/word_cloud
- [59] I. Gunhardsson, A. Svensson, C. Berterö, Documentation in palliative care: Nursing documentation in a palliative care unit—a pilot study 25 (1) 45–51, publisher: SAGE Publications Inc. doi:10.1177/1049909107307381.
URL <https://doi.org/10.1177/1049909107307381>
- [60] M. C. Broderick, A. Coffey, Person-centred care in nursing documentation 8 (4) 309–318, eprint: <https://onlinelibrary.wiley.com/doi/pdf/10.1111/opn.12012>. doi:10.1111/opn.12012.
URL <https://onlinelibrary.wiley.com/doi/abs/10.1111/opn.12012>
- [61] W.-C. Su, K. Dufendach, D. T. Wu, Assessing the readability of freely available ICU notes 2019 696–703.
URL <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC6568110/>
- [62] A. Celikyilmaz, E. Clark, J. Gao, Evaluation of text generation: A survey. arXiv:2006.14799 [cs], doi:10.48550/arXiv.2006.14799.
URL <http://arxiv.org/abs/2006.14799>