

Article

State-of-the-art in Open-domain Conversational AI: A Survey

Tosin Adewumi*, Foteini Liwicki and Marcus Liwicki

ML Group, EISLAB, Luleå University of Technology, Sweden; firstname.lastname@ltu.se

* Corresponding author

Abstract: We survey **SoTA** open-domain conversational **AI** models with the purpose of presenting the prevailing challenges that still exist to spur future research. In addition, we provide statistics on the gender of conversational **AI** in order to guide the ethics discussion surrounding the issue. Open-domain conversational **AI** are known to have several challenges, including bland responses and performance degradation when prompted with figurative language, among others. First, we provide some background by discussing some topics of interest in conversational **AI**. We then discuss the method applied to the two investigations carried out that make up this study. The first investigation involves a search for recent **SoTA** open-domain conversational **AI** models while the second involves the search for 100 conversational **AI** to assess their gender. Results of the survey show that progress has been made with recent **SoTA** conversational **AI**, but there are still persistent challenges that need to be solved, and the female gender is more common than the male for conversational **AI**. One main take-away is that hybrid models of conversational **AI** offer more advantages than any single architecture. The key contributions of this survey are 1) the identification of prevailing challenges in **SoTA** open-domain conversational **AI**, 2) the unusual discussion about open-domain conversational **AI** for low-resource languages, and 3) the discussion about the ethics surrounding the gender of conversational **AI**.

Keywords: conversational systems, chatbots, SotA

1. Introduction

There are different opinions as to the definition of **AI** but according to [1], it is any computerised system exhibiting behaviour commonly regarded as requiring intelligence. Conversational **AI**, therefore, is any system with the ability to mimic human-human intelligent conversations by communicating in natural language with users [2]. Conversational **AI**, sometimes called chatbots, may be designed for different purposes. These purposes could be for entertainment or solving specific tasks, such as plane ticket booking (task-based). When the purpose is to have unrestrained conversations about, possibly, many topics, then such **AI** is called open-domain conversational **AI**. ELIZA, by [3], is the first acclaimed conversational **AI** (or system). Its conversations with humans demonstrated how therapeutic its responses could be. Staff of [3] reportedly became engrossed with the program during interactions and possibly had private conversations with it [2].

Modern **SoTA** open-domain conversational **AI** aim to achieve better performance than what was experienced with ELIZA. There are many aspects and challenges to building such **SoTA** systems. Therefore, the primary objective of this survey is to investigate some of the recent **SoTA** open-domain conversational systems and identify specific challenges that still exist that should be surmounted to achieve "human" performance in the "imitation game", as described by [4]. As a result of this objective, this survey will identify some of the ways of evaluating open-domain conversational **AI**, including the use of automatic metrics and human evaluation. This work differs from previous surveys on conversational **AI** or related topic in that it presents discussion around the ethics of gender of conversational **AI** with compelling statistics and discusses the uncommon topic of conversational **AI** for low-resource languages. Our approach surveys some of the most representative work in recent years.

The key contributions of this paper are a) the identification of existing challenges to be overcome in **SoTA** open-domain conversational **AI**, b) the uncommon discussion



Citation: Title. *Preprints* 2022, 1, 0.
https://doi.org/

Publisher's Note: MDPI stays neutral with regard to jurisdictional claims in published maps and institutional affiliations.



Copyright: © 2022 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (https://creativecommons.org/licenses/by/4.0/).

about open-domain conversational AI for low-resource languages, and c) a compelling discussion about ethical issues surrounding the gender of conversational AI. The rest of the paper is organized as follows. The Background Section (2) presents brief details about some topics in conversational AI; the Benefits of Conversational AI Section (3) highlights some of the benefits that motivate research in conversational AI; the Methods Section (4) describes the details of the approach for the two investigations carried out in this survey; two Results of the Survey Sections (5 & 6) then follow with details of the outcome of the methods; thereafter, the Existing Challenges Section (7) shares the prevailing challenges to obtaining "human" performance; Open-domain Conversational AI for Low-resource Languages Section (8) discusses this critical challenge and some of the attempts at solving it; the Related Work Section (9) highlights previous related reviews, the Conclusion Section (11) summarizes the study after the limitations are given in the Limitation Section.

2. Background

Open-domain conversational AI may be designed as a simple rule-based template system or may involve complex artificial neural network (ANN) architectures. Indeed, six approaches are possible: (1) rule-based method, (2) reinforcement learning (RL) that uses rewards to train a policy, (3) adversarial networks that utilize a discriminator and a generator, (4) retrieval-based method that searches from a candidate pool and selects a proper candidate, (5) generation-based method that generates a response word by word based on conditioning, and (6) hybrid method [2,5,6]. Certain modern systems are still designed in the rule-based style that was used for ELIZA [2]. The ANN models are usually trained on large datasets to generate responses, hence, they are data-intensive. The data-driven approach is more suitable for open-domain conversational AI [2]. Such systems learn inductively from large datasets involving many turns in conversations. A turn (or utterance) in a conversation is each single contribution from a speaker [2,7]. The data may be from written conversations, such as the MultiWOZ [8], transcripts of human-human spoken conversations, such as the Gothenburg Dialogue Corpus (GDC) [9], crowdsourced conversations, such as the EmpatheticDialogues [10], and social media conversations like Familjeliv¹ or Reddit² [11,12]. As already acknowledged that the amount of data needed for training deep ML models is usually large, they are normally first pretrained on large, unstructured text or conversations before being finetuned on specific conversational data.

2.1. Retrieval & Generation approaches

Two common ways that data-driven conversational AI produce turns as response are Information Retrieval (IR) and generation [2]. In IR, the system fetches information from some fitting corpus or online, given a dialogue context. Incorporating ranking and retrieval capabilities provides additional possibilities. If C is the training set of conversations, given a context c , the objective is to retrieve an appropriate turn r as the response. Similarity is used as the scoring metric and the highest scoring turn in C gets selected from a potential set. This can be achieved with different IR methods and choosing the response with the highest cosine similarity with c [2]. This is given in Equation 1. In an encoder-encoder architecture, for example, one could train the first encoder to encode the query while the second encoder encodes the candidate response and the score is the dot product between the two vectors from both encoders. In the generation method, a language model or an encoder-decoder is used for response generation, given a dialogue context. As shown in Equation 2, each token of the response (r_t) of the encoder-decoder model is generated by conditioning on the encoding of the query (q) and all the previous responses ($r_{t-1}...r_1$), where w is a word in the vocabulary V . Given the benefit of these two methods, it may be easy to see the advantage of using the hybrid of the two for conversational AI.

¹ familjeliv.se

² reddit.com

$$response(c, C) =_{r \in C} \frac{c \cdot r}{|c| |r|} \quad (1)$$

$$r_t =_{w \in V} P(w|q, r_{t-1} \dots r_1) \quad (2)$$

2.2. Evaluation

Although there are a number of metrics for NLP systems [13–15] different metrics may be suitable for different systems, depending on the characteristics of the system. For example, the goals of task-based systems are different from those of open-domain conversational systems, so they may not use the same evaluation metrics. Human evaluation is the *gold standard* in the evaluation of open-domain conversational AI, though it is subjective [16]. It is both time-intensive and laborious. As a result of this, automatic metrics serve as proxies for estimating performance though they may not correlate very well with human evaluation [14,17,18]. For example, IR systems may use F1, precision, and recall [13]. Furthermore, metrics used in NLG tasks, like Machine Translation (MT), such as the BLEU or ROUGE, are sometimes used to evaluate conversational systems [16] but they are also discouraged because they do not correlate well with human judgment [2,19]. They do not take syntactic or lexical variation into consideration [15]. Dependency-based evaluation metrics, however, allow for such variation in evaluation. Perplexity is commonly used for evaluation and has been shown to correlate with a human evaluation metric called Sensibleness and Specificity Average (SSA) [5]. It measures how well a model predicts the data of the test set, thereby estimating how accurately it expects the words that will be said next [5]. It corresponds to the effective size of the vocabulary [13] and smaller values show that a model fits the data better. Very low perplexity, however, has been shown to suggest such text may have low diversity and unnecessary repetition [20].

Two methods for human evaluation of open-domain conversational AI are observer and participant evaluation [2]. Observer evaluation involves reading and scoring a transcript of human-chatbot conversation while participant evaluation interacts directly with the conversational AI [2]. In the process, the system may be evaluated for different qualities, such as humanness (or human-likeness), fluency, making sense, engagingness, interestingness, avoiding repetition, and more. The Likert scale is usually provided for grading these various qualities. The others are comparison of diversity and how fitting responses are to the given contexts. Human evaluation is usually modeled to resemble the Turing test (or the imitation game).

2.3. The Turing test

Modern human evaluation is generally designed like the Turing test. The Turing test is the indistinguishability test. This is when a human is not able to distinguish if the responses are from another human or a machine in what is called the imitation game [4]. The proposed imitation game, by [4], involves a man, a woman, and an interrogator of either sex, who is in a separate room from the man and the woman. The goal of the interrogator is to determine who is the woman and who is the man, and he does this by directing questions to the man and the woman, which are answered in some written format. The man tries to trick the interrogator into believing he's a woman while the woman tries to convince the interrogator she's a woman. When a machine replaces the man, the aim is to find out if the interrogator will decide wrongly as often as when it was played with a man [4]. The formulation of the imitation game, by [4], does not precisely match modern versions of the test [21].

An early version of the test was applied to PARRY, a chatbot designed by [22] to imitate aggressive emotions. Most psychiatrists couldn't distinguish between transcripts of real paranoids and PARRY [2,23]. This example of PARRY may be viewed as an edge case, given that the comparison was not made with rational human beings but paranoids [24]. A limited version of the test was introduced in 1991, alongside its unrestricted version, in what is called the Loebner Prize competition [24]. Every year, since then, prizes have been

awarded to conversational [AI](#) that pass the restricted version in the competitions [25]. This competition has its share of criticisms, including the view that it is rewarding tricks instead of furthering the course of [AI](#) [24,26]. As a result of this, [26] recommended an alternative approach, whereby the competition will involve a different award methodology that is based on a different set of assessment and done on an occasional basis.

2.4. Characteristics of Human Conversations

Humans converse using speech and other gestures that may include facial expressions, usually called body language, thereby making human conversations complex [2]. Similar gestures may be employed when writing conversations. Such gestures may be clarification questions or the mimicking of sound (*onomatopoeia*). In human conversations, one speaker may have the conversational initiative, i.e., the speaker directs the conversation. This is typical in an interview where the interviewer asking the questions directs the conversation. It is the style for [Question Answering \(QA\)](#) conversational [AI](#). In typical human-human conversations, the initiative shifts to and from different speakers. This kind of mixed (or rotating) initiative is harder to achieve in conversational systems [2]. Besides conversation initiative, below are additional characteristics of human conversations, according to [27].

- Usually, one speaker talks at a time.
- The turn order varies.
- The turn size varies.
- The length of a conversation is not known in advance.
- The number of speakers/parties may vary.
- Techniques for allocating turns may be used.

2.5. Ethics

Ethical issues are important in open-domain conversational [AI](#). And the perspective of deontological ethics views objectivity as being equally important [28–30]. Deontological ethics is a philosophy that emphasizes duty or responsibility over the outcome achieved in decision-making [31,32]. Responsible research in conversational [AI](#) requires compliance to ethical guidelines or regulations, such as the [General Data Protection Regulation \(GDPR\)](#), which is a regulation protecting persons with regards to their personal data [33]. Some of the ethical issues that are of concern in conversational [AI](#) are privacy, due to [personally identifiable information \(PII\)](#), toxic/hateful messages as a result of the training data and unwanted bias (racial, gender, or other forms) [2,16].

Some systems have been known to demean or abuse their users. It is also well known that machine learning systems reflect the biases and toxic content of the data they are trained on [2,34]. Privacy is another crucial ethical issue. Data containing [PII](#) may fall into the wrong hands and cause security threat to those concerned. It is important to have systems designed such that they are robust to such unsafe or harmful attacks. Attempts are being made with debiasing techniques to address some of these challenges [35]. Privacy concerns are also being addressed through anonymisation techniques [2,36]. Balancing the features of chatbots with ethical considerations can be a delicate and challenging work. For example, there is contention in some quarters whether using female voices in some technologies/devices is appropriate. Then again, one may wonder if there is anything harmful about that. This is because it seems to be widely accepted that the proportion of chatbots designed as “female” is larger than the those designed as “male”. In a survey of 1,375 chatbots, from automatically crawling chatbots.org, [37] found that most were female.

3. Benefits of Conversational [AI](#)

The apparent benefits inherent in open-domain conversational [AI](#) has spurred research in the field since the early days of ELIZA. These benefits have led to investments in conversational [AI](#) by many organizations, including Apple [2]. Some of the benefits include:

- Provision of therapeutic company, as was experienced with ELIZA.

- The provision of human psychiatric/psychological treatment on the basis of favorable behavior determined from experiments which are designed to modify input-output behaviour in models. This may be designed like PARRY [23].
- Provide support for users with disabilities, such as blindness [15].
- A channel for providing domain/world knowledge [15].
- The provision of educational content or information in a concise fashion [38].
- Automated machine-machine generation of quality data for low-resource languages [12].

4. Methods

We conduct two different investigations to make up this survey. Figure 1 depicts the methods for both investigations. The first addresses text-based, open-domain conversational AI in terms of architectures while the second addresses the ethical issues about the gender of such systems. The first involves online search on Google Scholar and regular Google Search, using the term "state-of-the-art open-domain dialogue systems". This returned 5,130 and 34,100,000 items in the results for Google Scholar and Google Search, respectively. We then sieve through the list of scientific papers (within the first ten pages because of time-constraint) to identify those that report SoTA results in the last five years (2017-2022) in order to give more attention to them. It is important to note that some Google Scholar results point to other databases, like ScienceDirect, arXiv, and IEEEExplore. The reason for also searching on regular Google Search is because it provides results that are not necessarily based on peer-reviewed publications but may be helpful in leading to peer-reviewed publications that may not have been immediately obvious on Scholar. We did not discriminate the papers based on the field of publication, as we are interested in as many SoTA open-domain conversational systems as possible, within the specified period. A second stage involves classifying, specifically, the SoTA open-domain conversational AI from the papers, based on their architecture. We also consider models that are pretrained on large text and may be adapted for conversational systems, such as the Text-to-Text Transfer Transformer (T5) [39], and autoregressive models because they easily follow the NLG framework. We do not consider models for which we did not find their scientific papers.

The second investigation, which addresses the ethical issues surrounding the gender of conversational AI, involves the survey of 100 chatbots. It is based on binary gender: male and female. The initial step was to search using the term "gender chatbot" on Google Scholar and note all chatbots identified in the scientific papers in the first ten pages of the results. Then, using the same term, the Scopus database was queried and it returned 20 links. The two sites resulted in 120 links, from which 59 conversational systems were identified. Since Facebook Messenger is linked to the largest social media platform, we chose this to provide another 20 chatbots. They are based on information provided by two websites on some of the best chatbots on the platform³. The sites were identified on Google by using the search term "Facebook Messenger best chatbots". They were selected based on the first to appear in the list. To make up part of the 100 conversational AI, 13 chatbots, which have won the Loebner prize in the past 20 years, are included in this survey. Finally, 8 popular conversational AI, which are also commercial, are included. These are Microsoft's XiaoIce and Cortana, Amazon's Alexa, Apple's Siri, Google Assistant, Watson Assistant, Ella, and Ethan by Accenture.

³ enterprisebotmanager.com/chatbot-examples
growthrocks.com/blog/7-messenger-chatbots

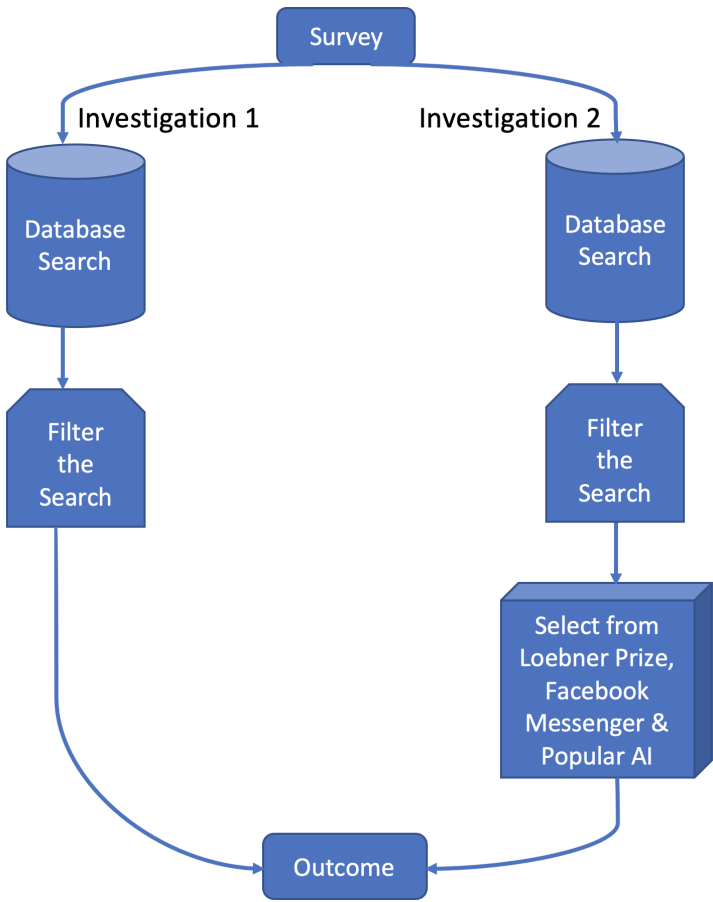


Figure 1. Method for both investigations in this study .

5. Results of Survey: Models

Review of the different scientific papers from the earlier method show that recent [SoTA](#) open-domain conversational [AI](#) models fall into one of the latter three approaches mentioned in Section 2: (a) retrieval-based, (b) generation-based, and (c) hybrid approaches. The models are BlenderBot 1 & 2, Meena, DLGNet, [Dialogue Generative Pre-trained Transformer \(DialoGPT\)](#), [Generative Pre-trained Transformer \(GPT\)-3](#) and 2 (RetGen), and [Text-to-Text Transfer Transformer \(T5\)](#).

5.1. BlenderBot 1 & 2

Some of the ingredients for the success of BlenderBot, as identified by [40], are empathy and personality, consistent persona, displaying knowledge, and engagingness. Three different parameter models are built for the variants: 90M, 2.7B, and 9.4B. The variants, which are all based on the Transformer, involve the latter three approaches: retrieval, generative, and a retrieve-and-refine combination of the earlier two. The generative architecture is a seq2seq model and uses Byte-Level [BPE](#) for tokenization. Human evaluation of multi-turn conversations, using ACUTE-Eval method, shows that its best model outperforms the previous [SoTA](#) on engagingness and humanness by using the Blended Skill Talk (BST) dataset [41]. They observed that models may give different results when different decoding algorithms are used, though the models may report the same perplexity in automatic metric. The more recent version of the set of models learns to generate an online search query from the internet based on the context and conditions on the results to generate a response, thereby employing the latest relevant information [42,43].

The seq2seq (or encoder-decoder) is an important standard architecture for BlenderBot and other conversational [AI](#) [43]. The Transformer, by [44], is often used as the underlying architecture for it, though the [Long Short Term Memory Network \(LSTM\)](#), by [45], may

also be used. Generally, the encoder-decoder conditions on the encoding of the prompts and responses up to the last time-step for it to generate the next token as response [2,5]. The sequence of tokens is run through the encoder stack's embedding layer, which then compresses it in the dense feature layer into fixed-length feature vector. A sequence of tokens is then produced by the decoder after they are passed from the encoder layer. The Softmax function is then used to normalized this, such that the token with the highest probability is the output.

5.2. Meena

Meena is presented by [5]. It is a multi-turn open-domain conversational AI seq2seq model that was trained end-to-end [46]. The underlying architecture of this seq2seq model is the Evolved Transformer (ET). It has 2.6B parameters and includes 1 ET encoder stack and 13 ET decoder stacks. Manual coordinate-descent search was used to determine the hyperparameters of the best Meena model. The data it was trained on is a filtered public domain corpus of social media conversations containing 40B tokens. Perplexity was used to automatically evaluate the model. It was also evaluated in multi-turn conversations using the human evaluation metric: Sensibleness and Specificity Average (SSA). This combines two essential aspects of a human-like chatbot: being specific and making sense.

5.3. DLGNet

DLGNet is presented by [47]. Its architecture is similar to GPT-2, being an autoregressive model. It is a multi-turn dialogue response generator that was evaluated, using the automatic metrics BLEU, ROUGE, and distinct n-gram, on the Movie Triples and closed-domain Ubuntu Dialogue datasets. It uses multiple layers of self-attention to map input sequences to output sequences. This it does by shifting the input sequence token one position to the right so that the model uses the previously generated token as additional input for the next token generation. Given a context, it models the joint distribution of the context and response, instead of modeling the conditional distribution. Two sizes were trained: a 117M-parameter model and the 345M-parameter model. The 117M-parameter model has 12 attention layers while the 345M-parameter model has 24 attention layers. The good performance of the model is attributed to the long-range transformer architecture, the injection of random informative paddings, and the use of BPE, which provided 100% coverage for Unicode texts and prevented the OOV problem.

5.4. DialoGPT 1 & 2 (RetGen)

DialoGPT was trained on Reddit conversations of 147M exchanges [16]. It is an autoregressive LM based on GPT-2. Its second version (RetGen) is a hybrid retrieval-augmented/grounded version. In single-turn conversations, it achieved SoTA in human evaluation and performance that is close to human in open-domain dialogues, besides achieving SoTA in automatic evaluation. The large model has 762M parameters with 36 Transformer layers; the medium model has 345M parameters with 24 layers; the small model has 117M parameters with 12 layers. A multiturn conversation session is framed as a long text in the model and the generation as language modeling. The model is easily adaptable to new dialogue datasets with few samples. The RetGen version of the model jointly trains a grounded generator and document retriever [48].

5.5. GPT-3 & GPT-2

GPT-3 is introduced by [49], being the largest size out of the eight models they created. It is a 175B-parameter autoregressive model that shares many of the qualities of the GPT-2 [50]. These include modified initialization, reversible tokenization, and pre-normalization. However, it uses alternating dense and locally banded sparse attention. Both GPT-3 and GPT-2 are trained on the CommonCrawl dataset, though different versions of it. GPT-3 achieves strong performance on many NLP datasets, including open-domain QA. In addition, zero-shot perplexity, for automatic metric, was calculated on the Penn Tree Bank

(PTB) dataset. Few-shot inference results reveal that it achieves strong performance on many tasks. Zero-shot transfer is based on providing text description of the task to be done during evaluation. It is different from one-shot or few-shot transfer, which are based on conditioning on 1 or k number of examples for the model in the form of context and completion. No weights are updated in any of the three cases at inference time and there's a major reduction of task-specific data that may be needed.

5.6. T5

T5 was introduced by [39]. It is an encoder-decoder Transformer architecture and has a multilingual version, mT5 [51]. It is trained on Colossal Clean Crawled Corpus (C4) and achieved SoTA on the SQuAD QA dataset, where it generates the answer token by token. A simplified version of layer normalization is used such that no additive bias is used, in contrast to the original Transformer. The self-attention of the decoder is a form of causal or autoregressive self-attention. All the tasks considered for the model are cast into a unified text-to-text format, in terms of input and output. This approach, despite its benefits, is known to suffer from prediction issues [52,53]. Maximum likelihood is the training objective for all the tasks and a task-prefix is specified in the input before the model is fed, in order to identify the task at hand. The base version of the model has about 220M parameters.

6. Results & Discussion of Survey: Ethics of Gender

Following the procedure mentioned in Section 4, each conversational AI's gender is determined by the designation given by the developer or cues such as avatar, voice or name, for cases where the developer did not identify the gender. These cues are based on general perception or stereotypes. We consider a conversational AI genderless if it is specifically stated by the reference or developer or nothing is mentioned about it and there are no cues to suggest gender. Overall, in the investigation of the 100 conversational AI, 37 (or 37%) are female, 20 are male, 40 are genderless, and 3 have both gender options. Figure 2 shows a bar graph with details of the results. Breaking down the data into 4 groups: journal-based, Loebner-winners, Facebook Messenger-based, and popular/commercial chatbots, we observe that female conversational AI always outnumber male conversational AI. The genderless category does not follow such a consistent trend in the groups. Out of the 59 chatbots mentioned in journal articles, 34% are female, 22% are male, 42% are genderless, and 2% have both gender options. 54% are female among the 13 chatbots in the Loebner-winners, 23% are male, 15% are genderless, and 8% have both options. Of the 20 chatbots from Facebook Messenger, 25% are female, 10% are male, 65% are genderless, and 0 offer both genders. Lastly, out of the 8 popular/commercial conversational AI, 62.5% are female, 25% are male, 0 is genderless, and 12.5% have both options.

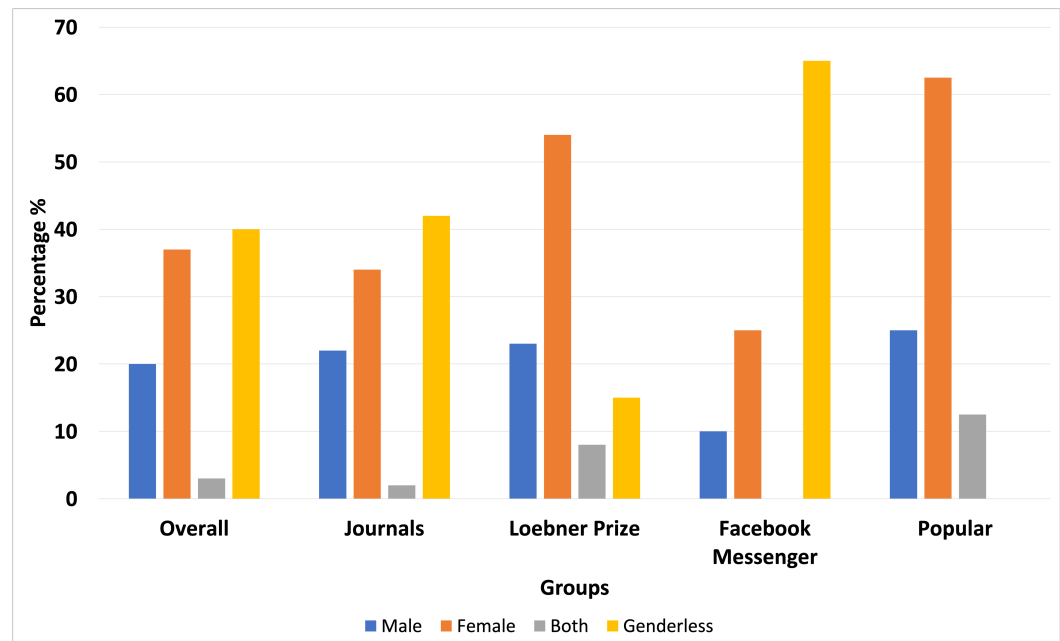


Figure 2. Bar chart of the gender of conversational AI for the Overall, Journal, Loebner prize, Facebook Messenger and Popular cases.

6.1. Discussion

The results agree with the popular assessment that female conversational AI are more predominant than the male ones. We do not know of the gender of the producers of these 100 conversational AI but it may be a safe assumption that most are male. This assessment has faced criticism from some interest groups, evidenced in a recent report by [54] that the fact that most conversational AI are female makes them the face of glitches resulting from the limitations of AI systems. Despite the criticism, there's the opinion that this phenomenon can be viewed from a vantage position for women. For example, they may be viewed as the acceptable face, persona or voice, as the case may be, of the planet. A comparison was made by [55] of a visually androgynous agent with both male and female agents and it was found that it suffered verbal abuse less than the female agent but more than the male agent. Does this suggest developers do away with female conversational AI altogether to protect the female gender or what is needed is a change in the attitude of users? Especially since previous research has shown that stereotypical agents, with regards to task, are often preferred by users [56]. Some researchers have argued that conversational AI having human-like characteristics, including gender, builds trust for users [57–59]. Furthermore, [60] observed that conversational AI that consider gender of users, among other cues, are potentially helpful for self-compassion of users. Noteworthy that there are those who consider the ungendered, robotic voice of AI uncomfortable and eerie and will, therefore, prefer a specific gender.

7. Existing Challenges of Open-domain Conversational AI

This survey has examined several SoTA open-domain conversational AI models. Despite their noticeable successes and the general progress, challenges still remain. The challenges contribute to the non-human-like utterances the conversational AI tend to have. These challenges also provide motivation for active research in NLP. For example, the basic seq2seq architecture is known for repetitive and dull responses [6]. One way of augmenting the architecture for refined responses is the use of IR techniques, like concatenation of retrieved sentences from Wikipedia to the conversation context [2]. Other shortcomings may be handled by switching the objective function to a mutual information objective or introducing the beam search decoding algorithm in order to achieve relatively more diverse responses [6]. Besides, GPT-3 is observed to lose coherence over really long passages, gives

contradictory utterances, and its size is so large that it's difficult to deploy. Collectively, some of the existing challenges are highlighted below. It is hoped that identifying these challenges will spur further research in these areas.

1. Poor coherence in sequence of text or across multiple turns of generated conversation [2,61].
2. Lack of utterance diversity [20].
3. Bland repetitive utterances [16,20].
4. Lack of empathetic responses from conversational systems [10].
5. Lack of memory to personalise user experiences.
6. Style inconsistency or lack of persona [5,16].
7. Multiple initiative coordination [2].
8. Poor inference and implicature during conversation.
9. Lack of world-knowledge.
10. Poor adaptation or responses to idioms or figurative language [18]
11. Hallucination of facts when generating responses [62].
12. Obsolete facts, which are frozen in the models' weights during at training .
13. Training requires a large amount of data [62].
14. Lack of common-sense reasoning [62].
15. Large models use so many parameters that make them complex and may impede transparency [62].
16. Lack of training data for low-resource languages [12,63]

8. Open-domain Conversational AI for Low-resource Languages

The last challenge mentioned in the earlier section is a prevailing issue for many languages around the world. Low-resource languages are natural languages with little or no digital data or resources [12,64]. This challenge has meant that so many languages are unrepresented in many deep ML models, as they usually require a lot of data for pretraining. Even Noteworthy, though, that multilingual versions of some of the models are being made with very limited data of the low-resource languages. They are, however, known to have relatively poor performance compared to models trained completely on the target language data [65–67] and only few languages are covered [12]. Approaches to mitigating this particular challenge involve human and automatic MT attempts [64] and efforts at exploiting cross-lingual transfer to build conversational AI capable of machine-machine conversations for automated data generation [12].

9. Related Work

In a recent survey, [68] reviewed advances in chatbots by using the common approach of acquiring scientific papers from search databases, based on certain search terms, and selecting a small subset from the lot for analysis, based on publications between 2007 and 2021. The databases they used are IEEE, ScienceDirect, Springer, Google Scholar, JSTOR, and arXiv. They analyzed rule-based and data-driven chatbots from the filtered collection of papers. Their distinction of rule-based chatbots as being different from AI chatbots may be disagreed with, especially when a more general definition of AI is given and since modern systems like Alexa have rule-based components [2]. Meanwhile, [69] reviewed learning towards conversational AI and in their survey classified conversational AI into three frameworks. They posit that a human-like conversation system should be both (1) informative and (2) controllable.

A systematic survey of recent advances in deep learning-based dialogue systems was conducted by [70], where the authors recognise that dialogue modelling is a complicated task because it involves many related NLP tasks, which are also required to be solved. They categorised dialogue systems by analysing them from two angles: model type and system type (including task-oriented and open-domain conversational systems). [71] also recognised that building open-domain conversational AI is a challenging task. They describe how, through the Alexa Prize, teams advanced the SoTA through context in dialog

models, using knowledge graphs for language understanding, and building statistical and hierarchical dialog managers, among other things.

10. Limitation

Although this work has presented recent SoTA open-domain conversational AI within the first ten pages of the search databases (Google Scholar & Google Search) that were used, we recognise that the time-constraint and restricted number of pages of results means there may have been some that were missed. This goes also for the second investigation on the gender of conversational AI. Furthermore, our approach did not survey all possible methods for conversational AI, though it identified all the major methods available.

11. Conclusion

In this survey of the SoTA open-domain conversational AI, we identified models that have pushed the envelope in recent times. It appears that hybrid models of conversational AI offer more advantages than any single architecture, based on the benefit of up-to-date responses and world knowledge. Besides discussing some of their successes or strengths, we focused on prevailing challenges that still exist and which need to be surmounted to achieve the type of desirable performance, typical of human-human conversations. The important challenge with conversational AI for low-resource languages is highlighted and the ongoing attempts at tackling it. The presentation of the discussion on the ethics of the gender of conversational AI gives a balanced perspective to the debate. We believe this survey will spur focused research in addressing some of the challenges identified, thereby enhancing the SoTA in open-domain conversational AI.

Author Contributions: “Conceptualization, Tosin Adewumi and Foteini Liwicki; methodology, Tosin Adewumi; investigation, Tosin Adewumi; writing—original draft preparation, Tosin Adewumi; writing—review and editing, Foteini Liwicki and Marcus Liwicki; supervision, Foteini Liwicki and Marcus Liwicki; All authors have read and agreed to the published version of the manuscript.”

Conflicts of Interest: “The authors declare no conflict of interest.”

References

1. States, U. Preparing for the future of artificial intelligence. **2016**.
2. Jurafsky, D.; Martin, J. *Speech and Language Processing: An Introduction to Natural Language Processing, Computational Linguistics, and Speech Recognition*; Dorling Kindersley Pvt, Limited, 2020.
3. Weizenbaum, J. A computer program for the study of natural language. *Fonte: Stanford*: <http://web.stanford.edu/class/linguist238/p36> **1969**.
4. Turing, A.M. Computing machinery and intelligence. *Mind* **1950**, *59*, 433–460.
5. Adiwardana, D.; Luong, M.T.; So, D.R.; Hall, J.; Fiedel, N.; Thoppilan, R.; Yang, Z.; Kulshreshtha, A.; Nemade, G.; Lu, Y.; et al. Towards a human-like open-domain chatbot. *arXiv preprint arXiv:2001.09977* **2020**. doi:10.48550/arXiv.2001.09977.
6. Chowdhary, K. Natural language processing. *Fundamentals of artificial intelligence* **2020**, pp. 603–649.
7. Schegloff, E.A. Sequencing in conversational openings 1. *American anthropologist* **1968**, *70*, 1075–1095.
8. Eric, M.; Goel, R.; Paul, S.; Sethi, A.; Agarwal, S.; Gao, S.; Kumar, A.; Goyal, A.; Ku, P.; Hakkani-Tur, D. MultiWOZ 2.1: A Consolidated Multi-Domain Dialogue Dataset with State Corrections and State Tracking Baselines. In *Proceedings of the 12th Language Resources and Evaluation Conference; European Language Resources Association: Marseille, France, 2020*; pp. 422–428.
9. Allwood, J.; Grönqvist, L.; Ahlsén, E.; Gunnarsson, M. Annotations and tools for an activity based spoken language corpus. In *Current and new directions in discourse and dialogue*; Springer, 2003; pp. 1–18. doi:10.1007/978-94-010-0019-2_1.
10. Rashkin, H.; Smith, E.M.; Li, M.; Boureau, Y.L. Towards Empathetic Open-domain Conversation Models: A New Benchmark and Dataset. In *Proceedings of the Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics; Association for Computational Linguistics: Florence, Italy, 2019*; pp. 5370–5381. doi:10.18653/v1/P19-1534.
11. Adewumi, T.; Brännvall, R.; Abid, N.; Pahlavan, M.; Sabry, S.S.; Liwicki, F.; Liwicki, M. Småprat: DialogPT for Natural Language Generation of Swedish Dialogue by Transfer Learning. In *Proceedings of the 5th Northern Lights Deep Learning Workshop, Tromsø, Norway. Septentrio Academic Publishing, 2022, Vol. 3*. doi:https://doi.org/10.7557/18.6231.
12. Adewumi, T.; Adeyemi, M.; Anuoluwapo, A.; Peters, B.; Buzaaba, H.; Samuel, O.; Rufai, A.M.; Ajibade, B.; Gwadabe, T.; Traore, M.M.K.; et al. Itàkùròso: Exploiting Cross-Lingual Transferability for Natural Language Generation of Dialogues in Low-Resource, African Languages **2022**. doi:10.48550/ARXIV.2204.08083.

13. Aggarwal, C.C.; Zhai, C. A survey of text classification algorithms. In *Mining text data*; Springer, 2012; pp. 163–222.
14. Gehrmann, S.; Adewumi, T.; Aggarwal, K.; Ammanamanchi, P.S.; Aremu, A.; Bosselut, A.; Chandu, K.R.; Clinciu, M.A.; Das, D.; Dhole, K.; et al. The GEM Benchmark: Natural Language Generation, its Evaluation and Metrics. In Proceedings of the Proceedings of the 1st Workshop on Natural Language Generation, Evaluation, and Metrics (GEM 2021); Association for Computational Linguistics: Online, 2021; pp. 96–120. doi:10.18653/v1/2021.gem-1.10.
15. Reiter, E. 20 Natural Language Generation. *The handbook of computational linguistics and natural language processing* **2010**, p. 574.
16. Zhang, Y.; Sun, S.; Galley, M.; Chen, Y.C.; Brockett, C.; Gao, X.; Gao, J.; Liu, J.; Dolan, B. DialoGPT: Large-Scale Generative Pre-training for Conversational Response Generation. In Proceedings of the Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics: System Demonstrations, 2020, pp. 270–278. doi:10.48550/arXiv.1911.00536.
17. Gangal, V.; Jhamtani, H.; Hovy, E.; Berg-Kirkpatrick, T. Improving Automated Evaluation of Open Domain Dialog via Diverse Reference Augmentation. In Proceedings of the Findings of the Association for Computational Linguistics: ACL-IJCNLP 2021; Association for Computational Linguistics: Online, 2021; pp. 4079–4090. doi:10.18653/v1/2021.findings-acl.357.
18. Jhamtani, H.; Gangal, V.; Hovy, E.; Berg-Kirkpatrick, T. Investigating Robustness of Dialog Models to Popular Figurative Language Constructs. In Proceedings of the Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing; Association for Computational Linguistics: Online and Punta Cana, Dominican Republic, 2021; pp. 7476–7485. doi:10.18653/v1/2021.emnlp-main.592.
19. Liu, C.W.; Lowe, R.; Serban, I.V.; Noseworthy, M.; Charlin, L.; Pineau, J. How not to evaluate your dialogue system: An empirical study of unsupervised evaluation metrics for dialogue response generation. *arXiv preprint arXiv:1603.08023* **2016**.
20. Holtzman, A.; Buys, J.; Du, L.; Forbes, M.; Choi, Y. The curious case of neural text degeneration. In Proceedings of the International Conference on Learning Representations, ICLR 2020, 2020.
21. Saygin, A.P.; Cicekli, I. Pragmatics in human-computer conversations. *Journal of Pragmatics* **2002**, *34*, 227–258.
22. Colby, K.M.; Hilf, F.D.; Weber, S.; Kraemer, H.C. Turing-like indistinguishability tests for the validation of a computer simulation of paranoid processes. *Artificial Intelligence* **1972**, *3*, 199–221.
23. Colby, K.M.; Weber, S.; Hilf, F.D. Artificial paranoia. *Artificial Intelligence* **1971**, *2*, 1–25.
24. Mauldin, M.L. Chatterbots, tinymuds, and the turing test: Entering the loebner prize competition. In Proceedings of the AAAI, 1994, Vol. 94, pp. 16–21.
25. Bradeško, L.; Mladenčić, D. A survey of chatbot systems through a loebner prize competition. In Proceedings of the Proceedings of Slovenian language technologies society eighth conference of language technologies. Institut Jožef Stefan Ljubljana, Slovenia, 2012, pp. 34–37.
26. Shieber, S.M. Lessons from a restricted Turing test. *arXiv preprint cmp-lg/9404002* **1994**.
27. Sacks, H.; Schegloff, E.A.; Jefferson, G. A simplest systematics for the organization of turn taking for conversation. In *Studies in the organization of conversational interaction*; Elsevier, 1978; pp. 7–55.
28. Adewumi, T.P.; Liwicki, F.; Liwicki, M. Conversational Systems in Machine Learning from the Point of View of the Philosophy of Science—Using Alime Chat and Related Studies. *Philosophies* **2019**, *4*, 41.
29. Javed, S.; Adewumi, T.P.; Liwicki, F.S.; Liwicki, M. Understanding the Role of Objectivity in Machine Learning and Research Evaluation. *Philosophies* **2021**, *6*, 22.
30. White, M.D. Immanuel kant. In *Handbook of economics and ethics*; Edward Elgar Publishing, 2009.
31. Alexander, L.; Moore, M. Deontological ethics **2007**.
32. Paquette, M.; Sommerfeldt, E.J.; Kent, M.L. Do the ends justify the means? Dialogue, development communication, and deontological ethics. *Public Relations Review* **2015**, *41*, 30–39.
33. Voigt, P.; Von dem Bussche, A. The eu general data protection regulation (gdpr). *A Practical Guide, 1st Ed., Cham: Springer International Publishing* **2017**, *10*, 10–5555.
34. Neff, G.; Nagy, P. Automation, algorithms, and politics | talking to Bots: Symbiotic agency and the case of Tay. *International Journal of Communication* **2016**, *10*, 17.
35. Dinan, E.; Fan, A.; Williams, A.; Urbanek, J.; Kiela, D.; Weston, J. Queens are Powerful too: Mitigating Gender Bias in Dialogue Generation. In Proceedings of the Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP); Association for Computational Linguistics: Online, 2020; pp. 8173–8188. doi:10.18653/v1/2020.emnlp-main.656.
36. Henderson, P.; Sinha, K.; Angelard-Gontier, N.; Ke, N.R.; Fried, G.; Lowe, R.; Pineau, J. Ethical challenges in data-driven dialogue systems. In Proceedings of the Proceedings of the 2018 AAAI/ACM Conference on AI, Ethics, and Society, 2018, pp. 123–129.
37. Maedche, A. Gender Bias in Chatbot Design. *Chatbot Research and Design* **2020**, p. 79.
38. Kerry, A.; Ellis, R.; Bull, S. Conversational agents in E-Learning. In Proceedings of the International conference on innovative techniques and applications of artificial intelligence. Springer, 2008, pp. 169–182.
39. Raffel, C.; Shazeer, N.; Roberts, A.; Lee, K.; Narang, S.; Matena, M.; Zhou, Y.; Li, W.; Liu, P.J. Exploring the Limits of Transfer Learning with a Unified Text-to-Text Transformer. *Journal of Machine Learning Research* **2020**, *21*, 1–67.
40. Roller, S.; Dinan, E.; Goyal, N.; Ju, D.; Williamson, M.; Liu, Y.; Xu, J.; Ott, M.; Shuster, K.; Smith, E.M.; et al. Recipes for building an open-domain chatbot. *arXiv preprint arXiv:2004.13637* **2020**.
41. Smith, E.M.; Williamson, M.; Shuster, K.; Weston, J.; Boureau, Y.L. Can You Put it All Together: Evaluating Conversational Agents' Ability to Blend Skills. In Proceedings of the Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics; Association for Computational Linguistics: Online, 2020; pp. 2021–2030. doi:10.18653/v1/2020.acl-main.183.

42. Komeili, M.; Shuster, K.; Weston, J. Internet-augmented dialogue generation. *arXiv preprint arXiv:2107.07566* **2021**.
43. Xu, J.; Szlam, A.; Weston, J. Beyond goldfish memory: Long-term open-domain conversation. *arXiv preprint arXiv:2107.07567* **2021**.
44. Vaswani, A.; Shazeer, N.; Parmar, N.; Uszkoreit, J.; Jones, L.; Gomez, A.N.; Kaiser, L.u.; Polosukhin, I. Attention is All you Need. In Proceedings of the Advances in Neural Information Processing Systems; Guyon, I.; Luxburg, U.V.; Bengio, S.; Wallach, H.; Fergus, R.; Vishwanathan, S.; Garnett, R., Eds. Curran Associates, Inc., 2017, Vol. 30.
45. Hochreiter, S.; Schmidhuber, J. Long short-term memory. *Neural computation* **1997**, 9, 1735–1780.
46. Bahdanau, D.; Cho, K.; Bengio, Y. Neural machine translation by jointly learning to align and translate. In Proceedings of the International Conference on Learning Representations, ICLR 2015, 2015. doi:10.48550/arXiv.1409.0473.
47. Olabiyi, O.; Mueller, E.T. Multiturn dialogue response generation with autoregressive transformer models. *arXiv preprint arXiv:1908.01841* **2019**.
48. Zhang, Y.; Sun, S.; Gao, X.; Fang, Y.; Brockett, C.; Galley, M.; Gao, J.; Dolan, B. Joint Retrieval and Generation Training for Grounded Text Generation. *arXiv preprint arXiv:2105.06597* **2021**.
49. Brown, T.; Mann, B.; Ryder, N.; Subbiah, M.; Kaplan, J.D.; Dhariwal, P.; Neelakantan, A.; Shyam, P.; Sastry, G.; Askell, A.; et al. Language models are few-shot learners. *Advances in neural information processing systems* **2020**, 33, 1877–1901.
50. Radford, A.; Wu, J.; Child, R.; Luan, D.; Amodei, D.; Sutskever, I. Language models are unsupervised multitask learners. *OpenAI blog* **2019**, 1, 9.
51. Xue, L.; Constant, N.; Roberts, A.; Kale, M.; Al-Rfou, R.; Siddhant, A.; Barua, A.; Raffel, C. mT5: A Massively Multilingual Pre-trained Text-to-Text Transformer. In Proceedings of the Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies; Association for Computational Linguistics: Online, 2021; pp. 483–498. doi:10.18653/v1/2021.naacl-main.41.
52. Adewumi, T.; Alkhaled, L.; Alkhaled, H.; Liwicki, F.; Liwicki, M. ML_LTU at SemEval-2022 Task 4: T5 Towards Identifying Patronizing and Condescending Language. *arXiv preprint arXiv:2204.07432* **2022**.
53. Sabry, S.S.; Adewumi, T.; Abid, N.; Kovacs, G.; Liwicki, F.; Liwicki, M. HaT5: Hate Language Identification using Text-to-Text Transfer Transformer. *arXiv preprint arXiv:2202.05690* **2022**.
54. West, M.; Kraut, R.; Ei Chew, H. I'd blush if I could: closing gender divides in digital skills through education **2019**.
55. Silvervarg, A.; Raukola, K.; Haake, M.; Gulz, A. The effect of visual gender on abuse in conversation with ECAs. In Proceedings of the International conference on intelligent virtual agents. Springer, 2012, pp. 153–160.
56. Forlizzi, J.; Zimmerman, J.; Mancuso, V.; Kwak, S. How interface agents affect interaction between humans and computers. In Proceedings of the Proceedings of the 2007 conference on Designing pleasurable products and interfaces, 2007, pp. 209–221.
57. Louwerse, M.M.; Graesser, A.C.; Lu, S.; Mitchell, H.H. Social cues in animated conversational agents. *Applied Cognitive Psychology: The Official Journal of the Society for Applied Research in Memory and Cognition* **2005**, 19, 693–704.
58. Muir, B.M. Trust between humans and machines, and the design of decision aids. *International journal of man-machine studies* **1987**, 27, 527–539.
59. Nass, C.I.; Brave, S. *Wired for speech: How voice activates and advances the human-computer relationship*; MIT press Cambridge, 2005.
60. Lee, M.; Ackermans, S.; Van As, N.; Chang, H.; Lucas, E.; IJsselsteijn, W. Caring for Vincent: a chatbot for self-compassion. In Proceedings of the Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems, 2019, pp. 1–13.
61. Welleck, S.; Kulikov, I.; Roller, S.; Dinan, E.; Cho, K.; Weston, J. Neural text generation with unlikelihood training. *arXiv preprint arXiv:1908.04319* **2019**.
62. Marcus, G. Deep learning: A critical appraisal. *arXiv preprint arXiv:1801.00631* **2018**.
63. Adewumi, T.P.; Liwicki, F.; Liwicki, M. The Challenge of Diacritics in Yoruba Embeddings. *arXiv preprint arXiv:2011.07605* **2020**.
64. Nekoto, W.; Marivate, V.; Matsila, T.; Fasubaa, T.; Fagbohunge, T.; Akinola, S.O.; Muhammad, S.; Kabongo Kabenamualu, S.; Osei, S.; Sackey, F.; et al. Participatory Research for Low-resourced Machine Translation: A Case Study in African Languages. In Proceedings of the Findings of the Association for Computational Linguistics: EMNLP 2020; Association for Computational Linguistics: Online, 2020; pp. 2144–2160. doi:10.18653/v1/2020.findings-emnlp.195.
65. Pfeiffer, J.; Rücklé, A.; Poth, C.; Kamath, A.; Vulić, I.; Ruder, S.; Cho, K.; Gurevych, I. AdapterHub: A Framework for Adapting Transformers. In Proceedings of the Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP 2020): Systems Demonstrations; Association for Computational Linguistics: Online, 2020; pp. 46–54.
66. Virtanen, A.; Kanerva, J.; Ilo, R.; Luoma, J.; Luotolahti, J.; Salakoski, T.; Ginter, F.; Pyysalo, S. Multilingual is not enough: BERT for Finnish. *arXiv preprint arXiv:1912.07076* **2019**.
67. Rönqvist, S.; Kanerva, J.; Salakoski, T.; Ginter, F. Is multilingual BERT fluent in language generation? *arXiv preprint arXiv:1910.03806* **2019**.
68. Caldarini, G.; Jaf, S.; McGarry, K. A Literature Survey of Recent Advances in Chatbots. *Information* **2022**, 13, 41.
69. Fu, T.; Gao, S.; Zhao, X.; Wen, J.r.; Yan, R. Learning towards conversational AI: A survey. *AI Open* **2022**, 3, 14–28.
70. Ni, J.; Young, T.; Pandelea, V.; Xue, F.; Adiga, V.; Cambria, E. Recent advances in deep learning based dialogue systems: A systematic survey. *arXiv preprint arXiv:2105.04387* **2021**.
71. Khatri, C.; Hedayatnia, B.; Venkatesh, A.; Nunn, J.; Pan, Y.; Liu, Q.; Song, H.; Gottardi, A.; Kwatra, S.; Pancholi, S.; et al. Advancing the state of the art in open domain dialog systems through the alexa prize. *arXiv preprint arXiv:1812.10757* **2018**.