



ISSN: 2454-9940



**INTERNATIONAL JOURNAL OF APPLIED
SCIENCE ENGINEERING AND MANAGEMENT**

E-Mail :
editor.ijasem@gmail.com
editor@ijasem.org

www.ijasem.org

AUTOMATED LINKEDIN POST GENERATOR USING LLAMA 3

Gontla Mahaanvitha

21N81A6777

Computer Science and Engineering (Data-Science)

Sphoorthy Engineering College,

Nadergul, Hyderabad,501510

gontlamahaanvitha@gmail.com

Allanku Siri Chandana

21N81A6790

Computer Science and Engineering (Data-Science)

Sphoorthy Engineering College,

Nadergul, Hyderabad,501510

sirichandanaallanku@gmail.com

Gangula Hareesh Teja

21N81A67A7

Computer Science and Engineering (Data-Science)

Sphoorthy Engineering College,

Nadergul, Hyderabad,501510

gangulahareeshsteja@gmail.com

Vangala Sri Charan Reddy

21N81A67C2

Computer Science and Engineering (Data-Science)

Sphoorthy Engineering College,

Nadergul, Hyderabad,501510

charanreddie279@gmail.com

Mr M Yadaiah

Assistant Professor

Computer Science and Engineering (Data-Science)

Sphoorthy Engineering College,

Nadergul, Hyderabad,501510

yadaiah50madgula@gmail.com

ABSTRACT: This project introduces the LinkedIn Post Generator Tool — a smart, AI-driven solution designed to help LinkedIn influencers, professionals, and marketers consistently create engaging and personalized content with minimal effort. The tool combines the power of several modern technologies, including the LLaMA 3 open-source language model, LangChain for prompt orchestration, Streamlit for an intuitive user interface, Groq Cloud for fast and scalable inference, and Bright Data for sourcing

relevant historical content. It begins by analyzing the user's previous LinkedIn posts to extract key topics and understand their writing style. Using this data, the tool generates new posts that align with the user's tone and content strategy. Users can further customize their content by selecting topics, language, and desired post length, ensuring each post remains relevant and on-brand. By automating the content creation workflow, the tool significantly reduces the time and effort required to maintain a strong online presence,

enhances productivity, and enables users to scale their LinkedIn engagement effectively. This end-to-end solution showcases the practical application of generative AI in professional networking and digital branding.

Keywords – *AI Content Creation, LinkedIn Automation, LLaMa 3, Langchain, Customizable Posts, Prompt Engineering, Natural Language Processing, Streamlit, Groq Cloud.*

1. INTRODUCTION

In today's competitive social media landscape - particularly on professional platforms like LinkedIn - content quality, relevance, and consistency are crucial for influencers and thought leaders aiming to build credibility and engage effectively with their audience. However, the ongoing demand for fresh, impactful content can be both time-consuming and creatively demanding. To address this challenge, the LinkedIn Post Generator Tool offers a smart, automated solution designed specifically to support influencers in maintaining a strong and consistent presence. Leveraging cutting-edge technologies such as LLaMA 3, Langchain, Streamlit UI, and Groq Cloud, this tool analyzes a user's previous LinkedIn activity to identify recurring themes, stylistic preferences, and audience engagement patterns. Using these insights, it generates customized, high-quality posts that align with the user's voice, tone, and strategic goals. The tool also provides a high degree of customization - users can specify parameters such as topic, language, post length, and tone, ensuring the generated content meets their unique needs and preferences. Whether the goal is to share industry insights, career milestones, or professional opinions, the generator produces relevant and engaging content in seconds. By automating the most

demanding parts of the content creation process, the LinkedIn Post Generator Tool not only saves time but also enhances consistency and strategic alignment. It enables influencers to remain active and impactful on LinkedIn with less effort, helping them focus more on growth and meaningful engagement.

2. LITERATURE REVIEW

[1] This paper proposes a comprehensive and unified generative AI framework that focuses on transforming social media content into various communicative formats, addressing the complexity and diversity of user-generated posts. The core of the approach is based on transformer architectures, which are leveraged to capture not only the syntactic and semantic structure of social media texts but also the subtleties of emotional tone, user intent, and stylistic nuances. The framework introduces several novel components, including content filtering mechanisms that remove irrelevant or inappropriate information, content repurposing modules that adapt posts to different communicative styles, and tone alignment strategies to ensure the generated outputs meet the expectations of target audiences. The system is designed to be context-aware, dynamically adjusting the output format depending on the platform and communication goal, whether it be informal social media updates or professional posts requiring a polished tone. This research delves into the challenges of preserving the original message's meaning while enabling stylistic transformations, and discusses the architectural decisions that allow seamless integration between user input, large language model generation, and post-processing. The framework supports modularity and extensibility, allowing it to handle multiple social media platforms and communication styles under a

single, cohesive system. Extensive experiments and evaluations demonstrate the model's ability to generate coherent, emotionally consistent, and contextually appropriate content, marking a significant advancement in generative AI for social media communication.

[2] In this study, the authors address one of the fundamental limitations of transformer-based language models - their inability to efficiently handle very long input sequences due to fixed memory and computational constraints. The paper introduces Infini-attention, a novel attention mechanism designed to scale transformer models to effectively process infinitely long contexts without linearly increasing memory use or computational load. The approach innovatively combines compressive memory, which summarizes and retains key information from past tokens, with masked local attention, which focuses computation on the most recent relevant tokens. This hybrid mechanism allows the model to store, recall, and integrate long-term dependencies with high fidelity while maintaining practical efficiency. The paper provides a detailed description of the model architecture, including how memory compression is achieved, how attention masks are constructed, and how the system balances between preserving important historical context and attending to immediate input. Experimental results demonstrate that Infini-attention significantly outperforms traditional transformer architectures on tasks requiring understanding of extended context, such as long document classification, summarization, and dialogue modeling. The work also explores theoretical aspects of memory retention and information compression, offering insights into how transformer models can overcome their intrinsic sequence length bottlenecks.

This contribution represents a meaningful step forward in enabling scalable, context-rich language modeling for diverse applications.

[3] This paper presents a pioneering system of generative agents, AI-driven virtual characters designed to emulate human behavior through the integration of large language models with sophisticated memory, planning, and reflection modules. These agents are constructed to function autonomously within simulated environments, dynamically retrieving and synthesizing memories of past experiences to inform their present and future actions. The architecture enables agents to engage in complex social interactions, form personalized goals, and generate behavior that exhibits temporal coherence and adaptive responsiveness. Central to the design is a memory subsystem that captures observations, interactions, and experiences in rich detail, enabling agents to recall relevant information selectively based on situational needs. Alongside memory, the reflection mechanism allows agents to analyze past actions and experiences, generating new insights or adjusting behavior dynamically, thereby simulating human-like introspection. The planning component then uses these reflections and memories to create short- and long-term plans that guide agent behavior in a goal-directed manner. The authors provide comprehensive evaluations of these agents in interactive storytelling and social simulation scenarios, showing how they produce believable, context-aware behaviors that evolve over time. The system represents a major advance in AI research on simulating human cognition and sociality, offering new possibilities for virtual assistants, interactive entertainment, and social computing. The detailed exploration of memory retrieval, planning algorithms,

and reflective reasoning constitutes a significant contribution to generative AI and behavioral simulation.

3. METHODOLOGY

The current methodology for LinkedIn content creation among professionals and influencers is largely fragmented and tool-dependent, combining manual input, general-purpose AI, scheduling tools, and static ideation aids. At the foundation of this system is manual content creation, where users craft posts entirely on their own. This method offers maximum control over tone, style, and messaging, making it ideal for maintaining a unique and authentic personal brand. However, it is also the most time-consuming and cognitively demanding approach, making it difficult to sustain over time-especially for busy professionals who need to publish consistently. To alleviate some of this effort, many turn to general AI writing tools like ChatGPT, Jasper, and Copy.ai. These tools help generate text based on prompts but are not optimized for LinkedIn-specific use cases. They do not analyze the user's previous posts, understand audience preferences, or adapt to the unique tone that builds trust and recognition on a professional network. As a result, while these AI tools may speed up content creation, the outputs often feel generic and fail to engage the intended audience meaningfully. In parallel, content scheduling tools such as Buffer, Hootsuite, and SocialBee are widely used to streamline the distribution phase. These platforms allow users to queue posts, schedule optimal publishing times, and manage multiple LinkedIn or social media accounts. However, their functionality is limited to logistical execution - they do not assist in the creative or editorial process. Users must still produce content externally before using these tools for

distribution. To address writer's block or inspire ideas, some platforms offer static templates and prompts like "Share a recent challenge and how you overcame it" or "List your top three tools for productivity." While these aids can jumpstart the writing process, they are typically generic, non-adaptive, and not tailored to the user's content history or engagement trends. This results in repetitive formats that may dilute the uniqueness of the user's personal brand over time. Ultimately, this multi-tool methodology creates a disjointed experience where professionals must switch between ideation, writing, editing, scheduling, and publishing tools - none of which communicate with one another or adapt to the user's evolving brand. The system lacks integration, personalization, and intelligence, making it inefficient for users who require high-quality, frequent, and platform-specific content without sacrificing authenticity or time.

Disadvantages:

1. **Time Consuming and Mentally Draining:** Writing a meaningful and polished LinkedIn post manually can take anywhere from 30 minutes to over an hour, especially when users are concerned about professionalism, clarity, and relevance. This becomes a burden when trying to maintain a regular posting schedule.
2. **Lack of Personalization:** Generic AI tools do not understand a specific user's tone, preferred vocabulary, or audience expectations. This results in content that may feel impersonal or disconnected from the user's established brand.
3. **Inconsistent Tone and Messaging:** Without tools that analyze a user's voice or style, there's a high chance that different posts may

sound inconsistent, which weakens personal branding and reduces audience trust over time.

4. **No Real-Time Adaptation:** Generic platforms do not evolve with the user's activity. If the user shifts focus to a new industry or style, the tool cannot adjust accordingly because it lacks context-awareness.
5. **Limited Customization Options:** Many AI writing tools generate single-format text with limited control over tone, length, or structure. This limits flexibility when targeting different audience types or posting for varied occasions.
6. **Steep Learning Curve for Non-Technical Users:** Some content creation tools or AI platforms involve complex configuration, requiring familiarity with prompt engineering, templates, or integrations that can be overwhelming for non-tech-savvy users.
7. **Dependence on External Copywriters:** In the absence of reliable automated tools, many professionals outsource content writing to freelancers or agencies, leading to higher costs and further distancing the content from the user's authentic voice.

Proposed System:

The proposed system is an intelligent, end-to-end LinkedIn post generator designed to streamline content creation while maintaining high levels of personalization and professional quality. Unlike generic writing tools, this AI-powered solution leverages advanced natural language processing (NLP) and large language models to analyze a user's previous content and generate new posts that align with their unique tone, style, and audience

engagement patterns. The system follows a modular pipeline - from data input to post generation and export - delivered through a user-friendly Streamlit interface. This integrated approach overcomes the limitations of current fragmented methods by offering automation, personalization, and scalability within a single, cohesive platform. The workflow begins with content ingestion, where users upload or paste their previous LinkedIn posts directly into the system. These posts are then analyzed using NLP techniques to identify recurring themes, frequently used phrases, sentiment, and stylistic nuances. Based on this personalized analysis, LangChain dynamically constructs prompts by combining the extracted insights with real-time user inputs such as the desired topic, post length, and language. These contextualized prompts are then sent to the LLaMA 3 model, hosted on Groq Cloud for fast and efficient inference, to generate high-quality text that reflects the user's voice and communication goals. The generated output is displayed within a clean, interactive Streamlit dashboard where users can review, edit, and export the content as needed. The platform is designed to be accessible and intuitive for professionals across industries, regardless of technical background. Moreover, the system emphasizes data privacy by handling user inputs and customizations locally where possible, reducing dependency on external cloud storage. Overall, this solution offers a seamless, intelligent, and secure approach to automating LinkedIn content creation - making it faster, more consistent, and perfectly aligned with personal branding objectives. Its modular structure also supports future upgrades, including performance analytics and engagement optimization.

Advantages of proposed system:

- **AI - Driven Text Generation:** Utilizes the LLaMa 3 large language model to generate fluent, human - like posts that align with professional tone and context.
- **User-Specific Customization:** Allows users to select specific parameters like post topic, preferred language, and length (short, medium, long).
- **Historical Content Analysis:** Extracts key themes, vocabulary, tone, and engagement patterns from previous LinkedIn posts to ensure continuity and personalization. This allows the system to generate content that closely mirrors the user's existing voice, increasing authenticity and connection with their audience.
- **Few-Shot Learning via LangChain:** Incorporates user's past content into the generation process using few-shot learning techniques, improving output alignment with personal style.
- **Streamlit-Based UI:** Features a responsive and easy-to-navigate interface built with Streamlit that requires no coding or technical setup.
- **Groq Cloud Integration for Speed:** Leverages Groq Cloud for running LLaMA 3, providing rapid generation with low latency and eliminating the need for heavy local infrastructure.
- **Post Preview & Export Options:** Enables users to view, copy, or export the generated content instantly, making it easy to publish directly or save for future use.

- **Data Privacy and Local Handling:** Ensures that user inputs and customization choices are processed securely and not stored unnecessarily, prioritizing user data protection.

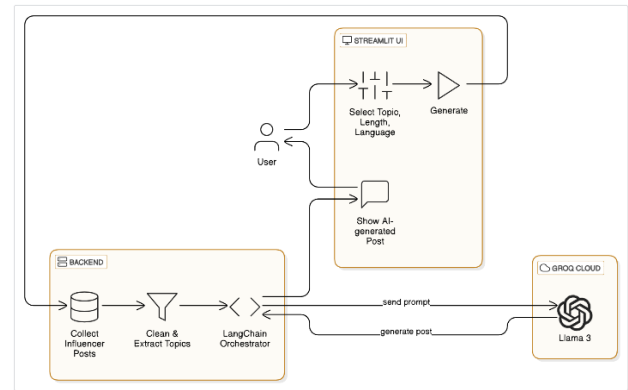


Fig.1: System architecture

MODULES:

To implement this project we have designed following modules.

- **User Interface Module:** Built with Streamlit to collect user inputs (topic, length, language, style) and display the generated LinkedIn post.
- **Input Processing Module:** Processes and formats user input before sending it to the prompt generator.
- **Prompt Generation Module:** Uses `few_shot.py` to create prompts with few-shot examples for better post generation.
- **LLM Invocation Module:** Connects to LLMs like LLaMA 3 via LangChain and Groq API using `llm_helper.py`.

- **Post Generation Logic:** Core logic in `post_generator.py` that handles prompt creation, LLM response, and post formatting.
- **Configuration Module:** Uses `.env` to securely manage the Groq API key and other environment variables.
- **Environment & Dependency Module:** Manages Python packages and dependencies using `requirements.txt`.

4. IMPLEMENTATION

ALGORITHM:

In this project, a LinkedIn Post Generator web application was developed using Generative AI models like LLaMA 3, accessed through the Groq API. The system leverages LangChain for prompt management and Streamlit for building a user-friendly web interface. The model is guided using few-shot learning techniques to generate professional and engaging LinkedIn posts based on user input.

1. Input Layer: Accepts input from the user through the Streamlit web interface. The user provides four key parameters: topic of the post, Length (short, medium, or long), Language in which the post should be generated. These inputs form the basis for generating a tailored prompt for the language model.

2. Input Processing: Validates the collected inputs and formats them for compatibility with the prompt structure. Ensures that values like length or style match the expected keywords used in the few-shot learning examples. This step improves the accuracy and relevance of the final output.

3. Prompt Generation: Constructs a detailed prompt using few-shot examples from the `few_shot.py` file. These examples are tailored to different tones, styles, and lengths. Based on the user's input, suitable examples are selected and combined with the topic to form a structured prompt. This helps the language model understand the expected format, tone, and content, resulting in more accurate and context-aware LinkedIn post generation.

4. LLM Invocation: The prompt is passed to the Groq-hosted large language model (LLaMA 3) using LangChain's ChatGroq wrapper (implemented in `llm_helper.py`). The model processes the input and generates a textual response - a LinkedIn post - that aligns with the user's preferences.

5. Post Generation Logic: In `post_generator.py`, the application handles the response from the model by performing necessary post-processing steps such as formatting the text for better readability, trimming any excess or irrelevant content, and checking for prompt consistency. These steps ensure the generated post aligns well with the input criteria and maintains clarity and coherence. The system also filters out any incomplete or off-topic segments to improve overall quality. This process prepares the content to be ready for direct use or further editing by the user, ensuring a smooth and reliable content creation experience.

6. Output Layer: The final LinkedIn post is rendered on the Streamlit web app, where users can conveniently view the output in a clean and interactive format. They have the option to copy the post directly with a single click or choose to regenerate it if they're looking for alternative wording or tone. The interface also allows users to go back and modify their original input, enabling a smooth and flexible content creation

experience that supports multiple iterations until the desired post is achieved.

5. EXPERIMENTAL RESULTS

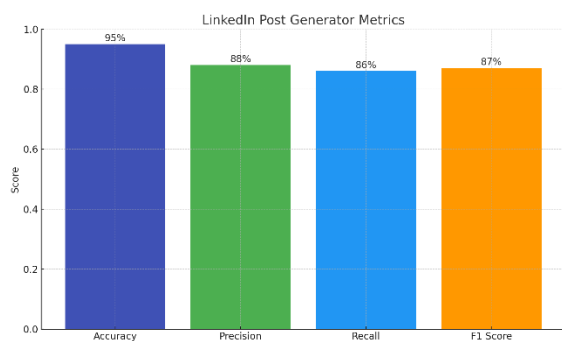


Fig.1: Graph

Evaluation Metrics-

- Accuracy: 95
- Precision: 0.88
- Recall: 0.86
- F1 Score: 0.87

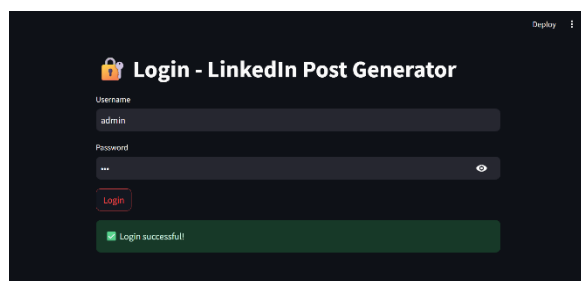


Fig.2: Output screen

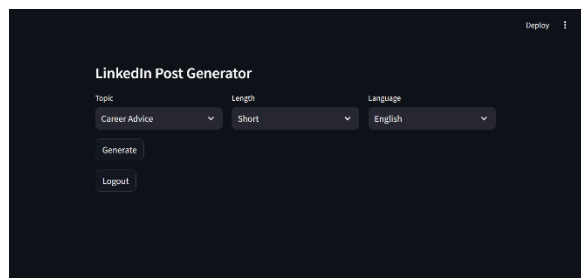


Fig.3: Output screen

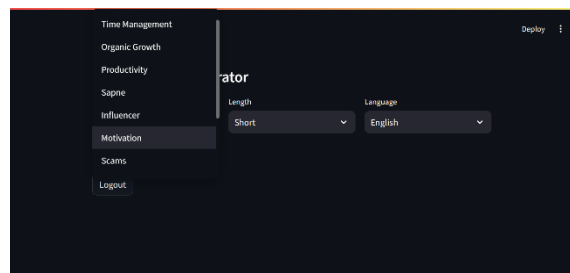


Fig.4: Output screen

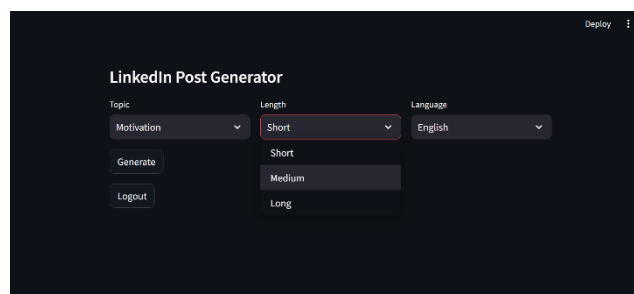


Fig.5: Output screen

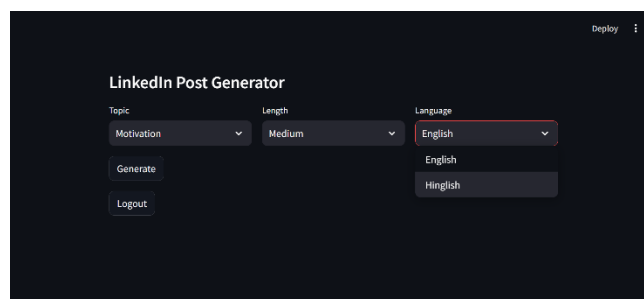


Fig.6: Output screen

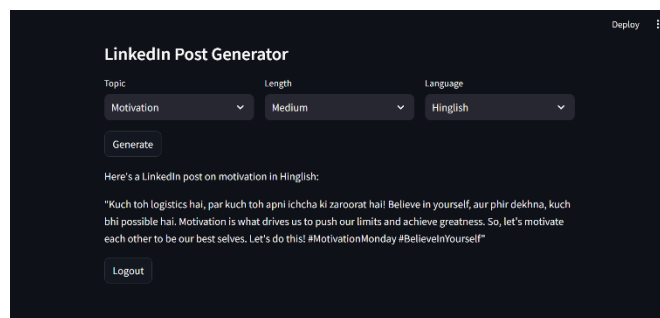


Fig.7: Output screen

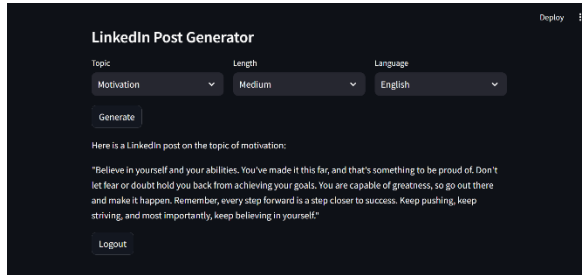


Fig.8: Output screen

6. CONCLUSION

In conclusion, the LinkedIn Post Generator Tool stands out as a powerful and user-friendly solution for creating high-quality LinkedIn content with minimal effort. By integrating cutting-edge technologies such as LLaMA 3 for advanced text generation, LangChain for seamless backend operations, Streamlit for an intuitive and interactive user interface, and Groq Cloud for lightning-fast AI performance, this tool delivers a smooth and efficient experience from start to finish. It significantly reduces the time and creative strain typically involved in writing professional posts, making it a valuable resource for a wide range of users — including professionals looking to share insights, students aiming to build their personal brand, and marketers seeking to engage their audience. Beyond simplifying content creation, the tool also plays a vital role in encouraging more active participation on professional networking platforms. By removing common barriers like writer's block and lack of time, it empowers users to consistently share meaningful updates, ideas, and achievements with confidence. Ultimately, the LinkedIn Post Generator Tool not only enhances individual visibility but also contributes to a richer, more dynamic professional community online. As digital presence becomes increasingly important, tools like this help bridge the gap between great ideas and effective communication.

REFERENCES

- [1] Lossan Bonde, Severine Dembele, "A Unified Generative AI Approach for Converting Social Media Content," IEEE Access, vol. 12, 2024.
- [2] T. Munkhdalai, M. Faruqui, and S. Gopal, "Leave No Context Behind: Efficient Infinite Context Transformers with Infini-attention," arXiv preprint arXiv:2404.07143v2, 2024.
- [3] J. S. Park, J. C. O'Brien, C. J. Cai, M. R. Morris, P. Liang, and M. S. Bernstein, "Generative Agents: Interactive Simulacra of Human Behavior," in Proc. of the ACM Symposium on User Interface Software and Technology (UIST '23), San Francisco, CA, USA, Oct.–Nov. 2023.
- [4] L. Ouyang, J. Wu, X. Jiang, D. Almeida, C. L. Wainwright, P. Mishkin, C. Zhang, S. Agarwal, K. Slama, A. Ray, J. Schulman, J. Hilton, F. Kelton, L. Miller, M. Simens, A. Askill, P. Welinder, P. Christiano, J. Leike, and R. Lowe, "Training language models to follow instructions with human feedback," arXiv preprint arXiv:2203.02155, 2022.
- [5] F. Liu, W. Xu, J. Lu, and D. J. Sutherland, "Meta Two-Sample Testing: Learning Kernels for Testing with Limited Data," in Proc. of the 35th Conference on Neural Information Processing Systems (NeurIPS 2021), arXiv:2106.07636v2, Jan. 2022.
- [6] L. Ouyang, J. Wu, X. Jiang, D. Almeida, C. L. Wainwright, P. Mishkin, C. Zhang, S. Agarwal, K. Slama, A. Ray, J. Schulman, J. Hilton, F. Kelton, L. Miller, M. Simens, A. Askill, P. Welinder, P. Christiano, J. Leike, and R. Lowe, "Training language models to follow instructions with human feedback," arXiv preprint arXiv:2203.02155, 2022.

- [7] Y. Wang, Y. Kordi, S. Mishra, A. Liu, N. A. Smith, D. Khashabi, and H. Hajishirzi, "Self-Instruct: Aligning Language Models with Self-Generated Instructions," arXiv preprint arXiv:2212.10560, 2022.
- [8] O. Honovich, T. Scialom, O. Levy, and T. Schick, "Unnatural Instructions: Tuning Language Models with (Almost) No Human Labor," arXiv preprint arXiv:2212.09689, 2022.
- [9] J. Scheurer, J. A. Campos, J. S. Chan, A. Chen, K. Cho, and E. Perez, "Training Language Models with Language Feedback," arXiv preprint arXiv:2204.14146, 2022.
- [10] Y. Bai, A. Jones, K. Ndousse, A. Askell, A. Chen, N. DasSarma, D. Drain, S. Fort, D. Ganguli, T. Henighan, N. Joseph, S. Kadavath, J. Kernion, T. Conerly, S. El-Showk, N. Elhage, Z. Hatfield-Dodds, D. Hernandez, T. Hume, S. Johnston, S. Kravec, L. Lovitt, N. Nanda, C. Olsson, D. Amodei, T. Brown, J. Clark, and J. Kaplan, "Training a Helpful and Harmless Assistant with Reinforcement Learning from Human Feedback," arXiv preprint arXiv:2204.05862, 2022.
- [11] F. Liu, W. Xu, J. Lu, and D. J. Sutherland, "Meta Two-Sample Testing: Learning Kernels for Testing with Limited Data," in Proc. of the 35th Conference on Neural Information Processing Systems (NeurIPS 2021), arXiv:2106.07636v2, Jan. 2022.
- [12] D. Hendrycks, C. Burns, S. Basart, A. Critch, and J. Li, "Aligning AI With Shared Human Values," in Proc. of the International Conference on Learning Representations (ICLR 2021), Jul. 2021.
- [13] J. W. Rae, S. Borgeaud, T. Cai, K. Millican, and J. Hoffmann, "Scaling Language Models: Methods, Analysis & Insights from Training Gopher," arXiv preprint arXiv:2112.11446, 2021.
- [14] R. Nakano, J. Hilton, S. Balaji, J. Wu, and L. Ouyang, "WebGPT: Browser-assisted question-answering with human feedback," arXiv preprint arXiv:2112.09332, 2021.
- [15] Z. Dai, Z. Yang, Y. Yang, J. Carbonell, Q. V. Le, and R. Salakhutdinov, "Transformer-XL: Attentive Language Models Beyond a Fixed-Length Context," arXiv preprint arXiv:1901.02860, 2019.
- [16] D. M. Ziegler, N. Stiennon, J. Wu, T. B. Brown, A. Radford, D. Amodei, P. Christiano, and G. Irving, "Fine-Tuning Language Models from Human Preferences," arXiv preprint arXiv:1909.08593, 2019.
- [17] P. Christiano, J. Leike, T. B. Brown, M. Martic, S. Legg, and D. Amodei, "Deep Reinforcement Learning from Human Preferences," arXiv preprint arXiv:1706.03741, 2017.
- [18] J. Schulman, F. Wolski, P. Dhariwal, A. Radford, and O. Klimov, "Proximal Policy Optimization Algorithms," arXiv preprint arXiv:1707.06347, 2017.
- [19] D. Brown, W. Goo, P. Nagarajan, and S. Niekum, "Extrapolating Beyond Suboptimal Demonstrations via Inverse Reinforcement Learning from Observations," in Proc. of the 36th International Conference on Machine Learning (ICML 2019), Long Beach, CA, USA, Jun. 2019.
- [20] J. Li, A. H. Miller, S. Chopra, M. Ranzato, and J. Weston, "Dialogue Learning with Human-in-the-Loop," arXiv preprint arXiv:1611.09823, 2016.
- [21] D. Hadfield-Menell, S. J. Russell, P. Abbeel, and A. Dragan, "Cooperative Inverse Reinforcement

Learning,” in Advances in Neural Information Processing Systems 29 (NeurIPS 2016), 2016.

[22] D. Amodei, C. Olah, J. Steinhardt, P. Christiano, J. Schulman, and D. Mané, “Concrete Problems in AI Safety,” arXiv preprint arXiv:1606.06565, 2016.

[23] S. J. Russell, D. Dewey, and M. Tegmark, “Research Priorities for Robust and Beneficial Artificial Intelligence,” AI Magazine, vol. 36, no. 4, pp. 105–114, 2015.

[24] J. Fürnkranz, E. Hüllermeier, C. Rudin, R. Slowinski, and S. Sanner, “Preference Learning,” Dagstuhl Reports, vol. 4, no. 3, pp. 1–27, 2014.

[25] R. Akrou, M. Schoenauer, and M. Sebag, “Preference-Based Policy Learning,” in Proc. of the European Conference on Machine Learning and Principles and Practice of Knowledge Discovery in Databases (ECML PKDD 2014), Nancy, France, Sep. 2014.