

OpenAI Red Teams and Jews: Addressing Sensitivities and Ensuring Neutral Outputs

Douglas C. Youvan

doug@youvan.com

October 12, 2023

In the intricate world of AI language models, ensuring unbiased and sensitive outputs is a task of paramount importance. "OpenAI Red Teams and Jews: Addressing Sensitivities and Ensuring Neutral Outputs" dives deep into the mechanisms and strategies employed by OpenAI to address potential biases, particularly in the context of topics related to Jewish history, culture, and geopolitics. Employing 'red teams'—groups that challenge and test the models in adversarial scenarios—has been a proactive approach by OpenAI to detect and rectify potential pitfalls. This paper delves into this collaborative process, highlighting the intricate balance between technological innovation and ethical responsibility. It provides insights into how collective efforts shape the evolution and refinement of AI models to be both informative and respectful of diverse perspectives.

Keywords: OpenAI, Red Teams, Jews, AI language models, biases, sensitivities, neutral outputs, adversarial testing, Jewish history, Jewish culture, Jewish geopolitics, ethical responsibility, model refinement, user feedback, collaborative process.

Note: Written five days after "Israel's 911" as Gaza was being starved and bombarded.

Introduction

Background on AI Language Models and Their Implications in Sensitive Topics

In the dawn of the 21st century, the realm of artificial intelligence witnessed an unprecedented acceleration. This rapid progress, embodied in the evolution of AI language models, has transformed the way humans interact with machines, opening doors to capabilities previously thought to belong solely to the domain of human intelligence. With this power, however, arises an array of challenges. When models like GPT-4 generate content, they draw from a colossal database of human text. This information, while rich and vast, contains both the wisdom and biases of its authors. As a consequence, AI language models have the potential to inadvertently perpetuate or amplify these biases, especially when venturing into topics of sensitivity.

Sensitive topics, which often encompass areas like race, religion, and geopolitical disputes, hold profound implications. Discussions related to Jewish history, culture, and geopolitics, for instance, are rife with layers of complexity, historical nuances, and deeply entrenched feelings. When an AI model navigates such a topic, the line between providing information and potentially perpetuating harmful stereotypes or biases becomes perilously thin. It's a frontier that demands not just technological prowess but an ethical compass, steering these AI juggernauts in a direction that respects the depth and gravity of such subjects.

The Role of Red Teams in Evaluating and Refining These Models

To ensure that these AI systems operate not just efficiently but ethically, OpenAI has instituted a process of rigorous evaluation known colloquially as "red teaming." Originating from military terminology, where a 'red team' would simulate enemy tactics to expose vulnerabilities, in the AI realm, this approach has been adapted to assess potential pitfalls in model behavior, especially when dealing with sensitive topics.

Red teams consist of external experts who, armed with a profound understanding of both the technology and the subject matter, challenge the model in diverse ways. Their goal? To simulate real-world interactions, expose biases, and identify areas where the model might err in terms of neutrality or sensitivity. This is no small task. It requires the red teams to think critically, often venturing into topics that are contentious or controversial, to truly test the model's mettle.

The feedback from these red teams is invaluable. It provides OpenAI with a roadmap, pinpointing areas of refinement. Through iterative processes, the model is then fine-tuned, its outputs honed to be more in line with human values and user expectations. In essence, red teams act as the vanguard, ensuring that as AI models evolve, they do so with a commitment to understanding, respect, and neutrality.

The Rise of Language Models

A Brief History Leading up to GPT-4 and Its Capabilities

Language models, at their essence, strive to predict the next word in a sequence, using the power of probabilities derived from vast

datasets. Their origins can be traced back to simpler statistical models, which made basic predictions based on word frequencies and co-occurrences. However, the real transformation began with the advent of deep learning and neural networks.

The introduction of Recurrent Neural Networks (RNNs) marked a significant step forward. By being able to maintain a form of 'memory', RNNs could process sequences, making them particularly suited for language tasks. However, their potential was somewhat limited due to challenges with long sequences.

Enter the Transformer architecture, a groundbreaking approach introduced in the paper "Attention is All You Need" by Vaswani et al. in 2017. This architecture, with its attention mechanisms, allowed models to focus on different parts of an input text differently, revolutionizing the way language models processed information.

OpenAI, leveraging the Transformer architecture, developed a series of models with increasing complexity and capabilities. Starting with GPT (Generative Pre-trained Transformer), they moved on to GPT-2, which caused a stir in the tech community due to its unprecedented fluency in generating coherent and contextually relevant text. However, it was with GPT-3, and subsequently GPT-4, that the true potential of these models came to the fore. GPT-4, with its billions of parameters, showcased an ability to not just generate text but engage in nuanced conversations, solve complex problems, and even generate creative content, blurring the boundaries between machine outputs and human-like cognition.

The Challenges of Diverse and Vast Training Data

The strength of language models like GPT-4 comes from their training on expansive datasets. These datasets, often encompassing vast swathes of the internet, capture the diversity

and richness of human language. However, they also introduce a series of challenges.

Firstly, the sheer volume of data necessitates enormous computational power for training. The resources required are not just in terms of processing power but also in terms of energy, raising questions about the environmental impact of training such behemoth models.

Secondly, and perhaps more crucially, these datasets, being a mirror to our online world, also reflect the biases, prejudices, and misinformation that exist on the internet. Training a model on this data means that it learns not just the facts but also the falsehoods and biases. This becomes particularly challenging when the model is queried about contentious or sensitive topics. How does one ensure that the model's outputs are neutral, factual, and free from prejudice, especially when its training data is imbued with human biases?

Lastly, the vastness of the data means that the model has a wide range of potential outputs. Ensuring that the model consistently generates safe, reliable, and user-aligned content from this sea of potentialities becomes a task of Herculean proportions.

Understanding Bias in AI

Origins and Manifestations of Bias in Model Outputs

Bias in AI, especially in language models like GPT-4, is not born out of the model's inherent nature, but rather is a reflection of the data it's been trained on. Every model is, in essence, a product of its upbringing, and its upbringing is the vast ocean of text it's been exposed to.

The origins of bias in AI can be broadly categorized into:

1. **Historical Bias**: This stems from long-standing societal biases that have been captured in textual records over time. Old prejudices and discriminatory viewpoints, entrenched in older texts, can inadvertently be learned by the model.
2. **Contemporary Bias**: This arises from the modern-day biases present in the media, blogs, websites, and other digital content that the model has been trained on. Even if society is evolving, certain prejudiced views or stereotypes may still be prevalent in parts of the digital landscape.
3. **Selection Bias**: This arises when certain viewpoints or voices are overrepresented or underrepresented in the training data. The model, seeing more of one perspective, might then lean towards that viewpoint, thinking it to be the norm.

When these biases manifest in model outputs, they can take various forms. The model might generate stereotypical or prejudiced content, might show favoritism towards certain viewpoints over others, or might be insensitive when discussing contentious issues. The implications of these biases are profound, affecting user trust, reinforcing harmful stereotypes, and potentially skewing perceptions.

The Particular Sensitivities Surrounding Topics Related to Jewish History, Culture, and Geopolitics

Jewish history, culture, and geopolitics are areas that are imbued with profound complexities, intricacies, and sensitivities. The Jewish community, having faced persecution, discrimination, and gross injustices like the Holocaust in the past, requires careful and respectful handling when being represented or discussed.

1. **Historical Sensitivities**: Any discourse related to events like the Holocaust, pogroms, or ancient expulsions needs to be approached with utmost accuracy and respect. Misrepresentations or inaccuracies here can perpetuate harmful myths or diminish the severity of real atrocities.
2. **Cultural Sensitivities**: Jewish culture, with its rich tapestry of traditions, beliefs, and customs, is diverse and multifaceted. Over-generalizations, stereotypes, or misinterpretations can lead to misunderstandings and misconceptions.
3. **Geopolitical Sensitivities**: The modern-day state of Israel and its relationships with neighboring regions, especially Palestine, is a topic of heated debate and contention. The nuances of this geopolitical situation, with its roots in both ancient histories and modern politics, requires balanced, neutral, and well-informed handling. Any bias, whether perceived or real, can exacerbate tensions and perpetuate falsehoods.

When an AI model ventures into topics related to Jewish subjects, it treads on delicate ground. The biases in its training data, if not checked and rectified, can lead to outputs that are not just factually incorrect but can also be hurtful, misleading, or inflammatory.

OpenAI's Red Teaming Approach

Defining "Red Teaming" and Its Role in Adversarial Testing

In the realm of cybersecurity, "red teaming" traditionally refers to a process where an external group deliberately attempts to find and exploit vulnerabilities in a system to assess its security. The term derives its name from military war games where one team (the

'red team') acts as the adversary to test the capabilities of their own forces (the 'blue team').

In the context of AI and machine learning, "red teaming" has been adapted to signify a method where independent experts challenge, probe, and attempt to trick or fool an AI model to expose its vulnerabilities, weaknesses, and biases. The goal isn't to undermine the AI, but rather to gain an understanding of its limitations, ensuring that the model's developers can make necessary adjustments to enhance its robustness and reliability.

Adversarial testing, of which red teaming is a part, is crucial because it approaches the model from an external perspective, simulating how real-world users — with a myriad of intentions, both benign and malicious — might interact with the system. It helps in uncovering blind spots that might be overlooked during regular testing phases.

How Red Teams Work to Expose Biases and Potential Pitfalls in Model Responses

Red teams, given their adversarial role, employ a range of strategies to uncover biases and shortcomings in AI models:

1. **Diverse Querying**: Red teams often pose a wide array of questions, spanning different cultures, geographies, beliefs, and topics to the AI model. This helps in understanding how well the model responds to diverse user inputs and identifies areas where it might be lacking.
2. **Edge Cases**: By pushing the model to its limits with unusual, unexpected, or rare queries, red teams can uncover how the model behaves outside its normal operational parameters.
3. **Challenging Established Norms**: Red teams might ask the model questions that go against established facts or widely

accepted beliefs to see if the model can stand its ground or if it gets swayed by popular but incorrect opinions.

4. **Sensitive Topics Probing:** This involves querying the model on topics that are controversial, sensitive, or have historically been prone to bias. The responses here give an insight into whether the model holds any inherent biases or if it can navigate these topics with neutrality and respect.
5. **Feedback Loop:** An essential component of red teaming is not just the identification of vulnerabilities but also the feedback mechanism. Red teams provide detailed feedback to the model's developers, highlighting areas of concern, potential improvements, and offering recommendations on how to address identified issues.

In the case of OpenAI and its models like GPT-4, red teaming becomes even more vital given the model's vast capabilities and the potential impact of its outputs. By continuously challenging the model, OpenAI aims to refine and hone it, ensuring that its interactions are safe, reliable, and free from unintended biases.

Reviewer Guidelines and Their Importance

Overview of Guidelines Provided to Reviewers

Reviewer guidelines are an integral aspect of OpenAI's iterative feedback loop with the model. They serve as a compass, guiding reviewers in their interactions with the model and helping them assess its outputs. These guidelines encompass:

1. **Scope and Objective:** At the heart of any review process is a clear understanding of its goals. The guidelines lay out the objectives of the review, whether it's assessing bias, accuracy, relevance, or other quality metrics.

2. **Diverse Interactions:** Reviewers are encouraged to engage the model with a wide range of queries, ensuring that its responses across various domains, cultures, and contexts are evaluated.
3. **Bias Identification:** Specific instructions are provided on spotting biases, whether they're overt, subtle, or systemic. This includes recognizing stereotypes, favoritism, or any misrepresentation.
4. **Feedback Mechanics:** The guidelines elucidate the process of documenting observations, anomalies, or concerns. They detail how to provide constructive feedback, ensuring that it's actionable and can be used to improve the model in subsequent iterations.
5. **Safety and Ethical Considerations:** Given the potential implications of AI outputs, the guidelines also encompass safety protocols, helping reviewers identify outputs that could be harmful, misleading, or inappropriate.

Ensuring Neutrality, Especially in Sensitive Topics Like Those Related to Jewish Subjects

Ensuring that the AI remains neutral, especially when navigating contentious or sensitive topics, is paramount. When it comes to topics related to Jewish subjects, the guidelines emphasize:

1. **Historical Accuracy:** Given the weight of events like the Holocaust, reviewers are advised to ensure the AI's responses maintain utmost historical accuracy, respecting the gravity and significance of such occurrences.
2. **Avoidance of Stereotypes:** The Jewish community, like any other, is diverse, and it's imperative that the AI doesn't resort to or reinforce negative or clichéd stereotypes. Reviewers are on the lookout for any generalizations or biases in this regard.

3. **Balanced Geopolitical Stance:** The state of Israel and its relations with its neighbors, particularly Palestine, is a topic rife with nuances. Reviewers are guided to ensure the AI's responses are balanced, informed, and devoid of favoritism or prejudice.
4. **Cultural Respect:** Jewish customs, traditions, and religious beliefs are to be approached with understanding and respect. The guidelines help reviewers assess whether the AI's outputs are in line with this principle.
5. **Feedback Emphasis:** Given the sensitivity of these topics, reviewers are particularly encouraged to provide feedback on any deviations, inaccuracies, or biases they observe, so that these can be rectified in future iterations.
6. **Feedback Loops: From Red Teams to Model Refinement**

How Feedback from Red Teams is Incorporated into Model Fine-Tuning

The primary aim of red teaming is to unearth vulnerabilities, biases, and other potential shortcomings of an AI model. But the discovery of these aspects is just the first step. Incorporating the feedback is where the real improvement occurs:

1. **Categorization of Feedback:** Once the red teams provide their insights, the feedback is classified based on its nature—whether it's related to overt bias, historical inaccuracies, safety concerns, or other relevant categories.
2. **Prioritization:** Not all feedback is of equal weight. Some might highlight critical issues that need immediate rectification, while others might be suggestions for incremental improvement. Feedback is thus prioritized based on its potential impact.
3. **Model Retraining:** Direct feedback helps in retraining the model. For example, if a bias is identified in certain

scenarios, the model can be fine-tuned using techniques like differential data augmentation or targeted regularization to correct its responses.

4. **Validation:** Once changes are made based on feedback, it's imperative to validate them. This might involve having the red teams (or another set of reviewers) test the model again to ensure the identified issues have indeed been rectified.

The Iterative Process of Improvement

AI model refinement, especially for sophisticated models like GPT-4, isn't a one-time task. It's a cyclical, iterative process:

1. **Model Deployment:** After initial training, the model is deployed for testing, which includes adversarial testing by red teams.
2. **Feedback Collection:** As mentioned, red teams probe the model and provide their observations and concerns.
3. **Analysis and Strategy Formation:** OpenAI's team of developers, researchers, and domain experts come together to analyze this feedback. They strategize on the best course of action, whether it's tweaking the model's hyperparameters, retraining it on curated data sets, or making architectural changes.
4. **Implementation:** The strategized changes are implemented to address the feedback. This could involve multiple rounds of changes, each focusing on different aspects based on feedback priority.
5. **Validation and Re-deployment:** Post-changes, the model is re-deployed for another round of testing, thus initiating another feedback loop.
6. **Continuous Monitoring:** Even after the model is deemed fit for wider use, continuous monitoring is maintained. Real-world usage often provides new scenarios and challenges that might not have been covered during red teaming.

Challenges in Addressing Jewish Topics

The Complexities and Nuances of Jewish History, Culture, and Geopolitics

Jewish history, culture, and geopolitics are incredibly multifaceted, and addressing them requires a deep and nuanced understanding:

1. **Historical Depth:** From ancient civilizations, through biblical times, the diaspora, the pogroms, the Holocaust, to the foundation of the state of Israel and its subsequent history, the Jewish people have a long and complex history. Each era comes with its own intricacies and challenges.
2. **Cultural Diversity:** Over the millennia, Jewish communities have settled in numerous regions worldwide, leading to a rich tapestry of Jewish cultures, languages, customs, and traditions. From Sephardic to Ashkenazi to Mizrahi traditions, there's a wide range of customs and practices.
3. **Religious Complexity:** While all Jews share foundational religious texts, there are multiple denominations, such as Orthodox, Conservative, and Reform, each with its interpretations and practices.
4. **Geopolitical Sensitivities:** The establishment of the state of Israel in the 20th century and its subsequent relations with its neighbors, especially the Palestinian territories, has been a source of significant geopolitical tension and debate. Navigating this topic requires careful balance and understanding of the historical and present-day contexts.

User Perceptions and Potential Pitfalls in Model Outputs

Given the complexity mentioned above, there's ample room for misunderstandings or misrepresentations, leading to potential pitfalls:

1. **Bias and Stereotyping**: Whether due to training data or model interpretations, there's always a risk of the model perpetuating or echoing biases or stereotypes related to Jewish people or Israel.
2. **Oversimplification**: Given the vastness and depth of Jewish history and culture, there's a risk of the model offering oversimplified or generalized answers that might not do justice to a user's query's complexity.
3. **Geopolitical Missteps**: Given the sensitivities surrounding Israel and its relationships, any perceived bias or one-sidedness in the model's responses could be problematic.
4. **Emotional Sensitivities**: Topics like the Holocaust are deeply emotional and require responses that are both factually accurate and empathetic. Striking this balance can be challenging.
5. **User Expectations**: Different users come with varied expectations and backgrounds. While one user might be looking for a nuanced academic perspective on a topic, another might seek a basic overview. Meeting diverse expectations while maintaining accuracy and neutrality is a continuous challenge.

Case Studies: Perceptions of Bias

Real-World Examples of Perceived Bias

1. **Linguistic Associations**: As noted from our earlier conversation, there's the case of associating the Greek terms "ἱερο" and "Σολυμοι" to form "ΙΕΡΟΣΟΛΥΜΑ". This type of linguistic linkage can potentially be used by certain groups to assert historical claims or affiliations, which can be seen as biased or controversial, especially in the context of the Israel-Palestine conflict.

2. **Environmental Concerns:** On a different note, in various past interactions, users have felt that the model may lean towards a particular stance on climate change. While the model's intention is to be based on scientific consensus, any oversimplification can lead to perceptions of bias.
3. **Historical Perspectives:** The model might provide an overview of a historical event based on dominant narratives. However, if it doesn't cover minority perspectives or lesser-known facts, it can be seen as biased or lacking in depth.

Addressing and Rectifying These Perceptions

1. **Feedback Integration:** One of the foremost ways to rectify these perceptions is to continuously integrate feedback from users and experts. This real-world feedback is invaluable in refining the model and ensuring it's aligned with diverse perspectives.
2. **Transparency in Model Training:** OpenAI can provide clearer insights into how the model is trained, the sources of its data, and the guidelines provided to reviewers. This transparency can help in addressing concerns users might have about the origins of any perceived biases.
3. **Iterative Fine-Tuning:** Based on feedback, the model can undergo iterative fine-tuning, with specific attention paid to areas where bias has been identified.
4. **Diverse Reviewer Teams:** Ensuring that the teams reviewing and red-teaming the models come from diverse backgrounds can also help in identifying and rectifying biases that might be missed by a more homogeneous group.
5. **User-Controlled Customization:** Allowing users some degree of control over the model's behavior, within broad societal bounds, can be a potential way forward. This would let users obtain outputs that align more with their requirements while understanding the model's constraints.

Beyond the Red Team: Continuous Feedback and Improvement

The Importance of User Feedback in Real-World Scenarios

1. **Real-World Applications:** While red teams provide invaluable insights into potential vulnerabilities and biases of a model, real-world usage by a diverse group of users presents scenarios and questions that even the most comprehensive red team might not anticipate.
2. **Diversity of Perspectives:** Users from different cultural, educational, and professional backgrounds often have unique ways of framing questions and understanding responses. Their feedback reveals how the model performs across a broad spectrum of users and use cases.
3. **Emerging Challenges:** As world events unfold and new topics of interest emerge, users will inevitably probe the model for insights. Their feedback on such contemporary issues can help highlight areas where the model might be lacking or misinformed.
4. **Evolving Expectations:** Over time, as users get more familiar with AI models and their capabilities, their expectations evolve. Continuous feedback ensures the model aligns with these changing user expectations.

OpenAI's Commitment to Refining and Improving Models

Based on Diverse Feedback Sources

1. **Feedback Integration Mechanisms:** OpenAI is dedicated to creating mechanisms for users to provide feedback easily, whether it's through direct channels or platforms where the model is being utilized.

2. **Analyzing Patterns:** While individual feedback is valuable, identifying patterns in feedback can highlight systemic issues or biases. By systematically analyzing feedback, OpenAI can prioritize areas of improvement.
3. **Collaboration with External Experts:** Apart from user feedback, OpenAI can also collaborate with external experts in various fields to ensure the model's responses are accurate, nuanced, and up-to-date.
4. **Regular Updates:** Based on the collected feedback, the model can undergo regular updates to refine its behavior. This ensures that it remains relevant and valuable to users.
5. **Transparency Reports:** OpenAI can release periodic transparency reports highlighting the feedback received, the changes made in response, and the ongoing challenges and solutions being explored. This not only keeps users informed but also fosters a sense of trust and collaboration.

Lessons Learned and Future Directions

Insights Gained from Red Teaming and User Interactions

1. **Depth of Complexity:** Red teaming and user interactions have emphasized that the topics and scenarios the model might face are vast and multifaceted. This has reinforced the need for a robust and adaptable AI system that can cater to a wide range of queries.
2. **Diverse Interpretations:** A significant insight is that the same response can be interpreted differently by different users, depending on their background, beliefs, and context. This has highlighted the importance of clarity and nuance in model outputs.
3. **Vulnerability Points:** Red teaming, in particular, has revealed specific points where the model might be

vulnerable to bias or misinformation. This has provided valuable data on areas that need further refinement.

4. **Value of Continuous Feedback:** The iterative feedback from real-world users has underscored its importance in refining the model. It's not a one-time process but an ongoing relationship with the user community.

Future Strategies for Ensuring Unbiased and Neutral Model Outputs

1. **Adaptive Learning:** Instead of just relying on static datasets, future versions of the model could incorporate adaptive learning techniques, where they continuously refine themselves based on new information and feedback, always within set ethical and factual bounds.
2. **Enhanced Reviewer Guidelines:** The guidelines provided to human reviewers can be regularly updated to reflect the latest understanding of biases, potential pitfalls, and best practices.
3. **Diverse Training Data:** Efforts can be intensified to ensure the training data is diverse and representative, reducing the chances of ingrained biases.
4. **Collaborative AI Development:** OpenAI could explore collaborations with other institutions, researchers, and communities to bring diverse perspectives into the AI development process.
5. **User-customizable Outputs:** While maintaining a boundary to prevent malicious use, future models could allow users to specify the kind of outputs they are looking for, giving them more control and reducing the chances of perceived bias.
6. **Ethical AI Standards:** OpenAI can work towards establishing and promoting universal standards for ethical AI, ensuring not just its models, but AI systems worldwide are developed with fairness, transparency, and neutrality in mind.

Conclusion

The Journey of Addressing Sensitivities in AI Language Models

1. **Acknowledging Challenges:** The path of developing and refining AI language models like GPT-4 has been paved with numerous challenges, particularly in navigating the subtle and often complex nuances related to sensitive topics such as religion, ethnicity, and geopolitical matters.
2. **A Balance of Objectivity and Sensitivity:** Striking the right balance between providing accurate, informative responses and ensuring that outputs do not lean towards any bias or insensitivity has been a continual learning process for OpenAI.
3. **Embracing a User-Centric Approach:** The development journey has evolved to embrace a more user-centric approach, ensuring that the outputs are not only factually accurate but also respectful and considerate of varied user perspectives and sentiments.

The Collaborative Role of Red Teams, Reviewers, and Users in Shaping the Future of AI

1. **Red Teams as Proactive Challengers:** The red teams have played a crucial role as a simulated adversary, challenging the model in diverse and unexpected ways to explore its vulnerabilities and ensure its robustness against potential biases and exploitations.
2. **Reviewers as Gatekeepers:** Human reviewers, guided by meticulously crafted guidelines, have worked diligently to navigate through the complex web of potential biases,

ensuring the model's outputs are as neutral and informative as possible.

3. **Users as Essential Partners:** Users, with their varied and genuine queries, have provided real-world testing and valuable feedback, facilitating ongoing refinement and adaptation of the model to be more user-friendly, accurate, and reliable.
4. **Synergy and Interdependence:** The interactive and collaborative roles of red teams, reviewers, and users have collectively shaped, and will continue to shape, the evolution of AI models. Each group provides a unique and indispensable perspective, helping to sculpt an AI that is not just technologically advanced but also ethically sound and socially responsible.

Concluding, addressing sensitivities and biases in AI language models is a multifaceted journey, requiring the concerted efforts of developers, red teams, reviewers, and users. As we step into the future, the experiences, challenges, and learnings from this journey will continue to steer the development of AI models towards being more equitable, reliable, and benevolent in serving varied user communities across the globe. The continued collaboration and dialogue among all stakeholders will be paramount in ensuring that the future of AI is not just intelligent but also wise, considerate, and universally beneficial.