

Group Debate Project: Topics and Guidelines

1. Assignment: In teams of four, you will be taking a side on an issue in ethics and technology, and engaging in a formal debate with four of your classmates. Your team will also write a short essay surveying the topic, and explaining why you have taken the side that you have. Details here:

- (a) The written component: Your team will write a roughly two page informal report on your assigned topic (500-700 words), to be turned in by the beginning of class on the day of our debate. Late papers will not be accepted. Roughly, you will briefly introduce the issue, and then provide some moral reasons in favor of both sides, and then briefly explain why your team believes that your position on this issue is the correct one.* More specific instructions for your particular prompt can be found on the pages below.

* Note: For the purposes of this assignment, when brainstorming reasons for and against each position, the focus of your efforts should be on the *moral* reasons for and against; though you *may* also appeal to practical and/or legal reasons—especially in instances where you are able to make a case that these other sorts of reasons might be morally relevant.)

- (b) The in-class debate: Here is how the debate will be structured:

Opening statement: Your team will have 3-5 minutes to deliver your opening remarks, outlining your team’s position and your reasons in favor of it. (You may elect one spokesperson, or everyone to participate in this oral presentation, or whatever. It’s up to you.) Team ‘NO’, which goes second, may even wish to incorporate some replies to the things said by Team ‘YES’ in their opening remarks (though this is not required).

Response statement: After both teams have given their opening statements, you will be given one minute to brainstorm together about what you would like to say in response to your opponents. You will then deliver a 2-3 minute response statement (e.g., providing your reasons for believing that what the opposing group has just said is mistaken).

Informal discussion: Then, the two teams will be given a few minutes to informally raise specific questions for one another, before I open discussion to the entire class for Q&A.

Order of Events	Time Allotment
Team “YES” opening statement	3-5 minutes
Team “NO” opening statement	3-5 minutes
Brainstorm session	1 minute
Team “YES” response statement	2-3 minutes
Brainstorm session	1 minute
Team “NO” response statement	2-3 minutes
Informal Q&A between teams	5-7 minutes
Informal Q&A from the audience	~25 minutes

6. How to Begin: Begin by reading your particular prompt (on the pages below). I then encourage you to do some preliminary brainstorming and research on your own, by yourself. As you begin to familiarize yourself with the issue and read articles about it, ask yourself: What is *my* moral stance on this issue? *Why* do I believe this? Why do others disagree? What *reasons* do they have for their stance?

After you have all thought about the issue on your own, you should then arrange to meet up with your teammates to discuss and share your thoughts from your preliminary research with one another. From there, you can then do some further brainstorming as a group, discuss strategies for defeating your opponents, make decisions about how best to divide up the work, and so on.

2. Grading Rubric: When grading, I will look for the following (roughly, in order of importance):

In-Class Debate

- *Preparedness*

It should be apparent that each of your team members has come to class prepared to discuss the issue—not only in the prepared portions of your presentation (if any), but also in the thoughtfulness of your team’s responses to your opponent’s claims. (For example, even though your team is arguing for *one* particular side of an issue, when preparing together, you should brainstorm reasons in favor of *both* sides, so that you can anticipate what your opponents will likely say on debate day, and so that you can come prepared and ready to refute their claims.) Finally, I will also look for evidence of advance preparation based on how you respond to audience questions during Q&A. (For example, if they ask something about a concern that’s pretty common or standard when considering this issue, it should be apparent that your team has already anticipated such a question, and thought about it in advance.)

- *Thoughtfulness/Critical Thinking and Reasoning*

Though your grade will not be based upon whether not you (appear to) “win” the debate, I *will* be looking for evidence that your team has staked out a thoughtful, well-argued position, and has come supplied with clear and persuasive reasons in favor of it (as well as clear and persuasive reasons against your opponents’ main claims). This will require your team to do some careful, critical thinking together behind the scenes.

- *Civility*

It should go without saying that our discussion will remain civil and respectful. This means no insulting of classmates, or shouting at or over them, and it also means giving others the opportunity to share their own views. We will also strive whenever possible to keep our comments constructive and productive, with the goal of moral progress and learning as we work through these difficult issues together, in a group effort.

Written Essay

- *Follows Instructions*

Your essay should be clear, engaging, and well-organized, and it should be apparent that your team has a good handle on the reasons both for and against their position, and is able to clearly explain *why* your team has taken up the position that it has, and you reject your opponent’s position (and in a way that is at least somewhat plausible and convincing). Though note: The essay component is included here primarily because (a) it gives me some guarantee that you all have thought about this topic in advance, before debate day—much as the video does for our in-class discussion days; (b) this is a writing-intensive seminar, and you’re required to write at least 6,000 words over the course of this semester. When grading, my primary focus will be on your performance during the debate. (See below.)

Note: Your grade will also take into consideration peer assessments of your performance, submitted by your teammates. So, please do your best to contribute your fair share to your group’s success.

3. Specific Topics: Specific prompts for the two debate topics are found on the pages below.

Debate One (Monday, 11/17): De-Platforming and Social Media

The Issue: There seems to be an awful lot of division, hatred, and misinformation going around these days. Do we have a moral duty to actively make it more difficult for these sorts of voices to be heard? That is, do we have a duty to “de-platform” certain people or beliefs?

[‘[No-Platforming](#)’ is the practice of refusing to give a person, or group of people, or a set of beliefs, etc., any public platform from which to speak, or share, or spread their views. No-platforming might include, for example, deleting a tweet or suspending a social media account, denying a permit to a rally organizer, refusing to allow someone to either rent out or be invited to a public venue where they want to give a talk, or refusing to publish pieces which endorse, promote, or normalize some particular belief.]

Let’s focus on the form most relevant to technology ethics: De-platforming on social media.¹ Recent examples of this include removing content from [climate change deniers](#), presidential [election deniers](#), COVID [anti-vaxxers](#), COVID [conspiracy theorists](#), [white nationalists](#), and [QAnon](#) conspiracy theorists.

In politics: In January 2021, Twitter suspended president Donald Trump’s account for tweets related to the January 6th insurrection. ([here](#); also [here](#)). January 2022, they removed Congresswoman Marjorie Taylor Greene for repeatedly spreading COVID misinformation. ([here](#)) Note that president Trump signed an executive order in 2020, forbidding censorship on social media—an order which social media platforms essentially ignored, and which president Biden reversed in 2021. ([here](#))

In 2022, Elon Musk further fanned the flames of this controversy after buying Twitter and vowing ([by tweet](#)) to end censorship on that platform, in the name of protecting free speech. (Though some question whether speech on Musk’s Twitter (‘X’), really is all that free; [here](#).) He then restored many of the suspended accounts. In 2025, Mark Zuckerberg suspends the fact-checking programs on Facebook and Instagram. ([here](#))

Moral Question: Do social media sites have a moral obligation to refuse to provide a platform to certain beliefs or ideologies? Is it at least *permissible* for them to censor certain beliefs?

[Note 1: It is uncontroversial that things like racial hate speech, death threats, or calls to incite or organize domestic terrorist attacks should be prohibited – and we already have laws banning such speech. But, let’s focus on the sorts of statements which are not so obviously immediately and directly harmful – e.g., content promoting climate change denial, anti-vaccination, or the claim that the 2020 presidential election was stolen. Should these beliefs be censored as well? For the purposes of this assignment, you should focus your discussion on this sort of content, which is where most of the present controversy is focused.]

¹ Though the recent controversy mostly focuses on *social media*, be aware that the debate over de-platforming initially gained momentum regarding *non*-internet platforms, in the wake of the white nationalist rally in Charlottesville in the summer of 2017. E.g.:

- **Rallies:** In August and September of 2017, many cities denied permits to white supremacist and alt-right groups who had sought to hold rallies across the U.S. (e.g., [Berkeley](#), [Syracuse](#), and more). Many other cities *granted* the permits (e.g., [Boston](#)). Interestingly, even the organizers of the Charlottesville rally had initially been *denied* a permit, until [the ACLU advocated](#) for their freedom of speech, stating, “The ACLU of Virginia stands for the right to free expression for all, not just those whose opinions are in the mainstream or with whom the government agrees.” (a verdict they stood by even after the event; listen [here](#))
- **Public Talks:** In October of 2017, Richard Spencer (one of the main organizers of the Charlottesville rally and the guy from the [‘punch a Nazi’ video](#)) was initially denied a request to give a [talk at the University of Florida](#)—a decision that was later overturned after threats of a lawsuit which charged the university with violating free speech. Interestingly, around that same time, Texas A&M University successfully [cancelled a campus rally](#) (a ‘White Lives Matter’ rally) organized by Spencer.
- **Editorials:** In November of 2017, *The New York Times* published a [profile piece on a Nazi sympathizer](#) from Ohio. The article immediately drew fierce criticism for portraying the subject as just a regular guy like you and me—seemingly *normalizing* Nazi principles and beliefs. Critics demanded that the article be removed, but [the editor](#) and [the author](#) ultimately stood by the article.

[Note 2: The question is not, Should the government *legally require* such sites to censor such content. But rather, *morally* speaking, *should* those companies censor such content? Legality and morality are two separate and very different issues. Consider: It is morally wrong to lie to your friends, so you shouldn't do it. But, there probably shouldn't be a law against it!]

For further discussion, see [this news report](#) on the issue. See also [here](#), [here](#), [here](#), [here](#), [here](#), and [here](#). See Mark Zuckerberg's speech against this sort of censorship [here](#) (follow-up [in 2025](#)).

Teams

You have been divided into team 'YES' and team 'NO'. Team YES will provide the first introductory remarks, and will argue that yes, social media platforms *do* have a moral obligation to remove the sort of potentially harmful content in question (e.g., climate denialism, anti-vax claims, etc.). Team NO will argue that they *do not* have an obligation to remove it (though they should still decide what stance they want to take regarding whether it is at least *permissible* for them to do so).

Optional Additional Readings

If you want to see what professional philosophers are saying about this issue, you might also wish to check out the following articles, available on Blackboard under 'Additional Readings':

- Robert Simpson and Amia Srinivasan, "No Platforming" (2018)
- Neil Levy, "No-Platforming and Higher-Order Evidence" (2019)
- Aluizio Couto, "Rescuing Liberalism from Silencing" (2021)

Day Two (Wednesday, 11/19): Large Language Models

The Issue: Sophisticated A.I. is here, and it is quickly becoming a huge part of our lives. About 200 million people already use ChatGPT every single day, for instance, and 35 million people use Google Gemini daily. (And that's even not counting all of the other LLMs, like DeepSeek, Claude, Grok, Ernie Bot, Microsoft Copilot, Perplexity AI, and so on.) With an average of 4.5 daily sessions per ChatGPT user, and an average of 14 minutes per session, that means that 200 million people are spending over an hour a day on average (63 minutes) using ChatGPT. And these numbers are expected to continue increasing steeply. (In fact, by the time you are reading this, I'm sure they are already much higher.)

Moral Question: Is this a good thing? Will the (near) inevitable saturation of large language models and other forms of AI into our daily lives be an overall good thing, or a bad thing?

I really love this excerpt from our sex robot reading by John Danaher:

“take the case of the iPhone (or smartphones, more generally) and ask yourself a simple question: What was Apple thinking when they introduced this product back in 2007? It was an impressive bit of technology, poised to revolutionize the smartphone industry, and set to become nearly ubiquitous within a decade. The social consequences were to be dramatic. Looking back, some of those consequences have been positive: increased connectivity, increased knowledge, and increased day-to-day convenience. But a considerable number of the consequences have been quite negative: the assault on privacy, increased distractability, endless social noise. Were any of these possible consequences weighing on the mind of Steve Jobs when he stepped onstage to deliver his keynote on January 9, 2007? Some possibly were, but ... it's unlikely he allowed the negative to hold him back for more than a few milliseconds. It was a cool product and it was bound to be a big seller. That's all that mattered. But when you think about it, this attitude is pretty odd. The success of the iPhone and subsequent smartphones has given rise to one of the biggest social experiments in human history. The consequences of near-ubiquitous smartphone use were uncertain at the time. Why didn't we insist on Jobs giving it a good deal more thought and scrutiny? Imagine if instead of an iPhone he was launching a revolutionary new cancer drug? In that case, we would have insisted upon a decade of trials and experiments, with animal and human subjects, before it could be brought to market. Why are we so blasé about information technology as compared to medication?”

We're engaging in a similar global social experiment again now—arguably, a *much bigger* one. We really ought to spend some time thinking about the potential implications of A.I. and LLMs for humanity, while we're still in the relatively early stages of these technologies.

Pros: Many tout LLMs (and other A.I.) as perhaps the greatest thing humanity has ever created, or ever will create. If used correctly, it might **solve every problem** in the world—climate change ([here](#)), disease ([here](#)), poverty and world hunger ([here](#)), and so on. Obviously, it will increase efficiency and **productivity** too (e.g., [here](#)). It can also be a super-**democratizer**, giving *everyone* access to the ability to express themselves effectively, or creatively, when they otherwise could not have, and giving *everyone* access to personalized education ([here](#)), the best medical ([here](#)) and legal advice ([here](#)), and so on. With a sophisticated LLM at your fingertips, almost anything is possible.

Cons: Or is this naïve? Could the rise of A.I. and LLMs instead be an overall bad thing? Of course, there are the concerns about technological **unemployment**, and the **existential threat** of a possible robot apocalypse, and worries associated with the increase in A.I.'s making **morally significant decisions** for us—all of which we've either already discussed, or will discuss very soon. But there are other concerns too. E.g., some worry that over-reliance on A.I. and LLMs will lead to a **decline**

in our cognitive capacities for critical thinking, reasoning, creativity, reading comprehension, and so on (see [here](#) and [here](#)); and an increased sense of **loneliness** (see [here](#) and [here](#)). It may also perpetuate **biases** (as was likely discussed by the facial recognition video/discussion group); or **stifle progress** by recycling only past human ideas. It could also undermine our confidence in or ability to discern the **truth** (e.g., if we're suddenly inundated with falsehoods, deepfakes, etc. – e.g., consider how bad bots *already* comprise [37% of all internet traffic](#)); not to mention **environmental concerns** associated with increased usage of A.I. (see [here](#) and [here](#)).

For further discussion, start with [this outstanding talk](#) from philosopher, Shannon Vallor. (For a shorter interview with Professor Vallor covering some of the highlights, read [here](#).) See also [this fascinating interview](#) with Open AI CEO, Sam Altman. Other interesting videos [here](#) and [here](#).

Teams

You have been divided into team 'YES' and team 'NO'. Team YES will provide the first introductory remarks, and will argue that yes, the rise of AI and LLMs will be an overall good thing for humanity. Team NO will argue that it *will not* be an overall good thing.

Note: Please avoid, if possible, focusing on the issues that we already have or will soon discuss elsewhere in this course (e.g., A.I. decision-making, technological unemployment, existential risks, responsibility gaps, etc.).

I'd also like to see you focus as much as possible on the short term, and at the individual level, rather than the long term or the global level. For example, how are LLMs already affecting *your own* lives? Have they made things better for you personally, overall, or worse? (I'm genuinely curious to hear your thoughts about this.) What do you foresee A.I. doing for you and others in the next five years or so—and will this be good or bad? And then, as a secondary focus, we can consider the further future, which is far less certain and more speculative (e.g., visions of LLMs ending world hunger and solving climate change, etc.).

Optional Additional Readings

If you want to see what professional philosophers are saying about this issue, you might also wish to check out the following articles, available on Blackboard under 'Additional Readings':

- Hendrik Kempt, "The Tech-Ethics of Large Language Models" (2025)
- Niina Zuber & Jan Gogoll, "Vox Populi, Vox ChatGPT: Large Language Models, Education and Democracy" (2024)
- *Optional:* Andrew Peterson, "AI and the Problem of Knowledge Collapse" (2025)