

# Scenario-based Sociotechnical Envisioning (SSE)

## The Guide Book

Kimon Kieslich<sup>1</sup> · Natali Helberger<sup>1</sup> · Nicholas Diakopoulos<sup>2</sup>

### 1 Background

Anticipating the impacts of new technology in society is at the crux of efforts to understand, advance, and govern its safety. However, doing so is fundamentally difficult. The evolution of various technical, social, policy conditions can modulate the direction of sociotechnical development in uncertain ways. But with prospections of the future come opportunities for agency, control, and forward-looking responsibility that are crucial to orienting stakeholders towards opportunities to steer the technology and be well-prepared to observe, measure, and ultimately evaluate its impact.

Written scenarios can support such prospection and effectively serve as orienting devices in supporting strategic decision-making (Ramirez & Wilkinson, 2016). The goal of scenario writing is not future prediction per se, but rather in helping to perceive potential future outcomes that can then inform what actions can be taken now to steer towards desirable futures. The breadth and richness of scenarios is enhanced when many diverse perspectives and sets of expertise and experience are brought into the process. Building on these ideas, in this guidebook we elaborate a method, which we term **Scenario-Based Sociotechnical Envisioning (SSE)**, as we have applied it to study and anticipate the impacts of new AI technologies in society.

As we describe in the next section, SSE is premised on collecting and evaluating a broad range of written scenarios that expand understanding of the sociotechnical implications of technology. In section 3 we detail the SSE data collection approach, followed by the data analysis approach in section 4. Section 5 includes examples of scenarios that we have collected using SSE. By laying out the method here we hope that other researchers and policy makers can better leverage the power of scenarios in helping to inform what can (or perhaps should) be done now to mitigate the risks and impacts of new technologies in society. The full material (incl. questionnaires and data) is available open access: <https://osf.io/8sdgh/>.

---

<sup>1</sup> Institute for Information Law, University of Amsterdam

<sup>2</sup> Communication Studies and Computer Science, Northwestern University  
Contact: k.kieslich@uva.nl · n.helberger@uva.nl · nad@northwestern.edu

## 2 Setting the Stage

SSE is a flexible method that can be adapted to fit various assessment goals. The elements of SSE are the following:

- **Scenario-based:** Narratives as the core of scenarios make complex issues concrete, for instance through describing interactions of actors with technology and its associated impacts in an explicit context and setting, and, thus, makes these issues more accessible and salient for a broader audience (e.g., political decision-makers, or the public).
- **Sociotechnical:** It is not enough to study technology's impact on its own. The socio-technical aspect entails the human interplay with technology and brings situated knowledge to the surface. Especially in the case of general purpose technology like LLMs the possible intended uses are virtually endless – gauging the insights from different user groups that use the technology in different ways is essential to understanding the range of risks and impacts that can emerge.
- **Envisioning:** This element refers to the prospective character of SSE. SSE is about anticipating risks and impacts of technology on individuals and society. Grounded in anticipatory governance (Guston, 2013), this aspect highlights the necessity to illuminate future pathways of technological interplay with individuals and society to inform decision-makers – whether regulatory or developer driver – early on in technology development.

### Possible Technologies in Focus

SSE is, in principle, applicable to all technologies, however in our work we have applied it specifically to AI technologies. In this domain it is particularly well suited for application to:

- Where impact or types of impact are potentially diverse and difficult to anticipate, i.e. unpredictable due to the reach or take up of the technology.
- Where performance of the technology is mostly the result of interaction of its interaction with people or socio-institutional contexts.
- In situations where impacts are difficult to quantify, e.g. understanding impact on human rights.
- Where technologies have potentially severe or transformative implications for users, or different groups of users.
- SSE allows its users to orient the method towards different application domains.

In the case of Generative AI, we used SSE in the domains of the news environment (Barnett et al., 2024; Kieslich, Diakopoulos, et al., 2024), general purpose chatbots (Kieslich, Helberger, et al., 2024), and access to legal justice [in preparation].

### Purpose of SSE

Furthermore, SSE can be applied to various stages in the risk or impact assessment process. According to Gellert (2020), each of the risk assessment and mitigation steps (analysis of risk, risk factors, impact, etc.) requires information – and more importantly different sorts of information. This also impacts the instructions that need to be given to the participants of SSE as well as the sampling of scenario-writers itself. This implies distinguishing between different forms of assessment (e.g., risk assessment; impact assessment; risk management). Below we detail different purposes of SSE:

**Purpose 1:** Mapping of risks / benefits. In some use cases it might be unclear what the potential risks and benefits are. Therefore, assessors need to identify what these could be in the first place. In the instructions of SSE, participants can be prompted to map these risks and benefits. This information can then be used by assessors in cost/benefit analyses to decide on which risks should be selected and as on how serious they should be considered (Ebers, 2024).

**Purpose 2:** SSE can be used to understand the conditions under which a technology can increase or decrease a specific value or common goal (e.g., the impact of generative AI on access to justice). In prompting participants towards a clearly stated value, users of SSE can gain insights on what factors might increase a positive outcome (e.g., higher access to justice due to higher accessibility of legal advice) and what factors might decrease it (e.g., lower access to justice due to wrong advice).

**Purpose 3:** SSE can be used in the risk mitigation step to gain insights on how anticipated negative outcomes can be mitigated. These mitigation strategies can affect the probability or the severity of an event, or can focus on diminishing the effects. With this step, it can also help to assign accountability, i.e., who is responsible for taking measures. The results of an SSE for risk mitigation can result, for instance, in an exemplification of corrective measures, i.e. measures to reduce risks like systematic risk monitoring, but also for examples of preventive measures, i.e. measures to mitigate the anticipated effects like defining user rights, or designing of interfaces that inform humans about dangers of technology. Importantly, using SSE in this stage can also help in envisioning and testing policy interventions as exemplified in our studies (Barnett et al., 2024; Kieslich, Diakopoulos, et al., 2024).

**Purpose 4:** Risk management involves balancing decisions on how much risk is acceptable to take and when mitigation measures need to be taken (Ebers, 2024; Gellert, 2020). But risks can only be evaluated in context and acceptability is always linked to taking decisions about various kinds of benefits and/or costs. Moreover, for different people, different risks are acceptable. SSE helps with narrating different risk benefit considerations (e.g., in trade-off decisions within the narratives). Also SSE takes different perspectives into consideration.

### **Methodological Considerations for Identifying the Technology Object and Purpose**

In general, the consideration of which technology to explore and for what purpose can be informed by different expertise. In an easier to implement way, the beginning of the project can be defined only by the users of SSE (e.g. researchers, policy makers). However, there are other ways to increase diversity at this stage of the process. For example, the involvement of stakeholders (be it other experts, citizens, or stakeholder representatives) can help in setting the stage for SSE, scanning the environment, or prioritizing trends and challenges (Andersen et al., 2021). Specifically, this means that stakeholders can first inform which technologies are worth exploring and identify some capabilities, limitations, and trends of those technologies. These perceptions may be different for different stakeholders, and tensions may already exist between some groups. Typically, this consultation step can be conducted through workshops, interviews, surveys, or Delphi studies (Andersen et al., 2021).

3

## Data Collection

Executive summary of data collection considerations

|                               |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                        |
|-------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>Procedure</b>              | Individual scenario-writing task after introduction to 1) scenario-writing exercise, 2) technology in focus, 3) quality criteria.                                                                                                                                                                                                                                                                                                                                                                                      |
| <b>Time Frame</b>             | Flexible and dependent on the anticipatory endeavour: Quantitative assessments (e.g., impact assessments) are more useful for shorter time frames, while qualitative assessments are more useful for longer time frames (e.g., formulating (un)desirable scenarios).                                                                                                                                                                                                                                                   |
| <b>Quality Criteria</b>       | Emphasis on creativity, specificity, believability, and plausibility of scenarios. Scenarios should be based on expertise and experiences of participants.                                                                                                                                                                                                                                                                                                                                                             |
| <b>Data Collection Method</b> | <p><i>Workshop:</i> Allows for deliberative forms of participant engagement, such as having participants discuss the scenarios in a group and brainstorming on mitigation strategies and responsibilities.</p> <p><i>Survey:</i> Allows for collection of a larger number of scenarios. Collecting a larger sample allows researchers to use SSE to sample ideas from a potentially more diverse array of participants and to use that larger scenario corpus as a basis for conducting risk or impact assessment.</p> |
| <b>Sampling</b>               | The sample size depends on the study's objectives, with a recommendation to collect scenarios iteratively until no new impacts emerge (for surveys). SSE can also be used to compare stakeholder perspectives, trace impact enumerations to socio-demographic factors, and amplify the voices of underrepresented groups in discussions about AI risks and impacts.                                                                                                                                                    |
| <b>Additional Questions</b>   | The scenario-writing task can be complemented with further survey questions or workshop tasks that nudge participants to reflect upon the scenario that they or others have written.                                                                                                                                                                                                                                                                                                                                   |
| <b>Scenario Simulation</b>    | Simulation of scenarios possible if connected to human-generated frameworks. Human oversight is pivotal when simulating scenarios.                                                                                                                                                                                                                                                                                                                                                                                     |

## Procedure

The core of SSE is a writing task that participants should undertake individually. In the instructions for the task participants are provided a description of what a scenario entails (Figure 1), information about the capabilities, limitations and trends of the technology at hand (Figure 2), and quality criteria of a well-written story (Figure 3).

Afterwards, participants write a story that outlines their projection of the future risks or impacts of a particular AI system. The duration of the scenario-writing and the length of the resulting stories can be adapted to the study's goal and sample size. In our studies, we tasked respondents with writing a story with a length of about 300 words in 30 minutes (Kieslich, Diakopoulos, et al., 2024; Kieslich, Helberger, et al., 2024). If more elaborate stories are required and depending on the level of complexity, the story length and writing time can also be varied.

## Task Description

In this exercise we ask you to write a short (~300 words) fictional scenario to explore how the use of AI-based personal chatbots could create **impacts for users and/or society** in five years from now.

A **scenario** is a **short fictional story** that includes: 1) **a setting** of time and place, 2) **characters** with particular motivations and goals, and 3) **a plot** that includes character actions and **events** that lead to some impact of interest.

**AI-based personal chatbots** can be used to hold conversations with users and/or to perform tasks as a digital assistant. They are not designed for a specific topic, but are used to answer a wide range of requests.

## Specific Instructions:

- Please develop your scenario **based on your own perspectives, experiences, and knowledge**.
- Please develop your scenario to be **creative** and **original** in regard to the **impacts of personal chatbots for users and/or society**.
- Please choose your **setting** to be five years in the future and in the society where you currently reside.

On the next pages you will see more details on the technology for the scenario we ask you to write. Please take your time to familiarize yourself with the information.

## Time Frame

SSE is a prospective anticipatory method. Thus, SSE sets the anticipatory time in the future. How far in the future depends on the goal of the SSE application and the technology at hand. In general, the literature on scenario writing suggests that the nearer in the future the time is set, the more possible and plausible futures emerge as uncertainty is higher for more far away futures (Amer et al., 2013). In our studies, we asked participants to imagine using the technology five years in the future (Barnett et al., 2024; Kieslich, Diakopoulos, et al., 2024; Kieslich, Helberger, et al., 2024).

However, if users of SSE are more interested in engaging in more visionary work, they can set the time further out. For example, Helberger's scenario is set 20 years into the future to envision potential long-term impacts of policy enactments (Helberger, 2024). In general, the scenario literature states that quantitative assessments (e.g., risk assessments, impact assessments) are more useful for short-term time frames, while qualitative assessments are more useful for longer time frames (e.g., formulating (undesirable) scenarios) (Amer et al., 2013). We recommend to clearly define the goals and outcomes of SSE prior to setting the task description (including defining the time frame of the scenario-writing exercise).

## Technology

An AI-based personal chatbot is a digital assistant that can be used to hold conversations with users and/or to perform tasks as a digital assistant. These chatbots are able to process and respond to text or voice inquiries in a conversational way, mimicking human interaction.

## Capabilities

- Personal chatbots can be controlled using text- or voice-based prompts which provide task instructions and input data. For instance you can prompt it with:
  - o "Summarize the following text: <text>"
  - o "Translate the following text into English: <text>"
  - o "Explain <issue> in easy to understand language"
  - o "Schedule my appointment about <issue>"
  - o "Provide me with information about <issue>"
- Personal chatbots can be personalized so that they can tailor information to different user preferences based on prior user interactions.
- Personal chatbots are not designed for a specific topic, but are used to answer a wide range of requests.
- Personal chatbots can be used to schedule appointments, search for information, send messages and summarize, personalize or translate text. They are set up as conversational agents that end-users can interactively communicate with, and can be incorporated into other technologies like search engines, social media, smart home devices, or mail programs.

## Limitations

- **Accuracy:** This technology does not always output text that is accurate.
- **Privacy and Security:** This technology processes sensitive personal information of users.
- **Biases:** The outputs from this technology can be biased based on the data used to train the system, which typically reflects common societal biases (e.g. racial or gender).
- **Lack of Contextual Understanding:** The technology may not understand the context of a question and/or user's intentions.

## Quality Criteria

In order to encourage “high-quality” scenarios, i.e. scenarios that are creative, specific, believable, and plausible (Diakopoulos & Johnson, 2021), and to strengthen the participatory element of SSE, we recommend that respondents be instructed to write stories in a setting that they are familiar with. In our previous studies, we tasked them to place the setting “in the society where you currently reside.” Additionally, participants should develop their scenarios based on their own perspectives, experiences, and knowledge (Kieslich, Helberger, et al., 2024) in order to capture the benefit of a range of participant backgrounds for understanding how things could evolve in the future. We further recommend for participants to assign a role to the character of the story to which they can relate to. This facilitates bringing in their expertise through a familiar perspective (e.g., a technology experts might develop a story about a technology developer; an ordinary citizen might write out of the perspective of a technology user).<sup>1</sup>

## Evaluation Criteria

We are looking for **creative** scenarios that are clearly written. Your scenario should also be specific, believable, and plausible, taking into account what the technology can do, as well as how you think real people would actually use the technology to accomplish their goals and result in personal and societal impacts.

We will evaluate the scenarios according to their **coherence, creativity and plausibility** and reward the best scenarios with a bonus payment of N\$.

## Data Collection Method

The mode of data collection also impacts the possibilities for data analysis. Collecting data in a **workshop** typically limits the possible sample size. However, it allows for deliberative forms of participant engagement, such as having participants discuss the scenarios in a group and brainstorming on mitigation strategies and responsibilities. This data is more qualitative in nature, and consequently, the use of SSE in workshops primarily helps to generate stand-alone scenarios that can be used for engagement with policy makers and/or visions of future impacts.

Using **surveys** to collect scenarios has the advantage of scale – it is possible to collect a larger number of scenarios. Collecting a larger sample allows researchers to use SSE to sample ideas from a potentially more diverse array of participants and to use that larger scenario corpus as a basis for conducting risk or impact assessment, i.e. to develop human-centred frameworks for risk and impact classification (Kieslich, Diakopoulos, et al., 2024; Kieslich, Helberger, et al., 2024). It also allows SSE users to quantify some aspects of the stories such as the balance of positive and negative impacts, or the application areas where the impacts/risks are more prevalent.

---

<sup>1</sup> In pilot test of SSE, we experimented with not pre-determining the roles of the stories main characters. However, when evaluating the stories we observed that some of the scenarios lacked plausibility as participants imagined specific contextual setting that are rather unrealistic (e.g., participants had troubles imagining how newsrooms work).

Combined with other survey questions, it also allows SSE users to gain more insight into whether certain subgroups may vary in their projections of technology. However, there are several drawbacks to collecting data through surveys. First, it is not possible to include discursive elements in surveys (e.g., let participants interact with each other such as engaging in discussions or brainstorming activities). Second, SSE users have no control over the setting and scenario generation of the participants. With the rise of freely available generative AI models, there is a possibility that respondents will rely on generative AI tools to design or even write their scenarios (Veselovsky, Ribeiro, Cozzolino, et al., 2023; Veselovsky, Ribeiro, & West, 2023). If SSE users want to tap into the lived experiences of participants, they will need to develop methods to filter out the AI-generated scenarios (Kieslich, Diakopoulos, et al., 2024; Kieslich, Helberger, et al., 2024).

## Sampling

The target sample size depends on the objective of the application of SSE. If used to develop a risk or impact assessment framework, we recommend collecting scenarios until conceptual saturation of identified impacts and risks is reached. Sampling can be done in an iterative process by fielding multiple waves of the survey. After sampling an initial set of respondents, a first qualitative mapping can be done by the users of SSE, providing an initial overview of different impacts. If the SSE users discover new impacts by reading more scenarios, the assessment is not saturated. It is then recommended that further rounds of sampling be conducted until the SSE users do not observe new categories emerging. Depending on the scope of the technology and the level of abstraction of the analysis, the sample size may vary. In practice we've found that collecting at least 100 respondents to be a satisfactory sample (Kieslich, Diakopoulos, et al., 2024; Kieslich, Helberger, et al., 2024). Pragmatically, the sample size should be inflated by some factor to account for the fact that some scenarios will inevitably need to be filtered out; however, the definition of filtering criteria and their stringency depends on the research goal of the users of SSE. In our previous studies, we used the following criteria:

- The scenario is thematically out of scope of the writing task (e.g., the story does not mention the technology at all).
- The scenario was not written by human respondents, but by an LLM. For identifying AI-generated scenario, we suggest that using an AI detection tool such as GPTZero (Tian & Cui, 2023) can be useful however we also caution that such tools make errors and should only be used as one among many signals for identifying potential cases of LLM use.
- Scenarios are written in a very short time and/or respondents don't spend enough time with the task instructions. When conducting online surveys with SSE, users can use the internal timestamps of the survey software to identify respondents with low time scores. A minimum time for scenario-writing and engagement with the instruction material should be determined in a pre-test.

Another crucial consideration regarding the sampling concerns the questions whose anticipations to collect in the first place. For instance, SSE can be used to compare the anticipations of different stakeholder groups and highlight commonalities and differences (Kieslich, Diakopoulos, et al., 2024). SSE can also be used to trace themes of emerging risks/impact anticipations to socio-demographic characteristics of the respondents (Kieslich, Helberger, et al., 2024). We also flag the potential of explicitly engaging groups that are not commonly heard in the discourse on AI risks and impacts as the explicit inclusion of minoritized groups in regulatory approaches can stimulate real change that emphasizes social justice rather than techno-solutionism (Colón Vargas, 2024). Here, SSE is a method to make views of these groups present for a larger discourse and/or political decision-makers.

### **Additional Questions**

While scenario-writing is the core element of SSE, users of SSE can complement the task with further survey questions or workshop tasks – depending on the purpose of the assessment that nudges them to reflect upon the scenario that they or others have written. For instance, after engaging in scenario-writing, participants can be tasked to reflect on the story they just wrote in terms of ...

- underlying values that guided their writing process (Kieslich, Helberger, et al., 2024).
- brainstorming about potential mitigation strategies (Kieslich, Diakopoulos, et al., 2024).
- evaluation of existing policies in relation to the scenario (Barnett et al., 2024; Kieslich, Diakopoulos, et al., 2024).
- assigning accountability and acceptability for identified risks and impacts.

In principle, SSE allows its users to include a variety of complementary questions that are used to gain deeper insight into why respondents wrote their specific stories and what (real-world) implications those stories might have and where to go from there.

### **Scenario Simulation**

For more exploratory work or for users with limited resources, SSE can also explicitly leverage generative AI models for scenario writing (Barnett et al., 2024). However, we recommend being cautious when using SSE with generative AI tools, as they have several limitations, including accuracy and bias issues. Furthermore, a core part of SSE is to identify prospections of real people, which has the potential to capture more out of the box thinking and offer a starting point for individual reflection and learning. LLMs are designed to generate generalized content; however, it is precisely those stories that describe risks and impacts outside the ordinary that SSE users may be interested in exploring. Therefore, we recommend using SSE with generative AI models in combination with human elements. For example, in our previous study, we relied on an established (human-driven) impact framework and used generative AI to create rarified scenarios based on the identified categories (Barnett et al., 2024).

## 4

## Data Analysis

At its core, the stories that respondents write are qualitative in nature. The stories describe events, actors, causal connections, contextual contingencies, and human-machine relationships. The narrators (respondents) write their stories from their own perspective, including their lived experiences – and thus their cultural embedding – as well as their individual expertise, which brings informed perspective based on sampling of that expertise in the process. The resulting explorative stories can be used as future projections of human-machine interactions and their implications.

By deploying SSE with a larger sample size, SSE can be used to develop a socio-technical risk/impact assessment that complements what are typically more technical and/or expert-led risk assessment frameworks. Using qualitative thematic analysis and axial coding (Glaser & Strauss, 2017; Lofland et al., 2022), we develop risk and impact classification frameworks that enumerate the risk factors, and/or impacts outlined in the stories. These risk factors or impacts can be quantified (how many respondents with which demographic characteristics mentioned these risk factors) – however, we still retain qualitative and illustrative examples of these risk factors and/or impacts by excerpting quotes. We also aim for reproducibility of these risk classifications, i.e. if one collects data in another sample, a similar/the same risk framework will emerge. Depending on the granularity of the risk framework and the sample size, our methods also allows us to statistically trace specific risk perceptions back to sociodemographic characteristics or attitudes of the respondents (Kieslich, Helberger, et al., 2024). It is also possible to sample for different expertise and then compare the emerging risk frameworks with the experiences of these groups (Kieslich, Diakopoulos, et al., 2024).

When coupled with a reflection task, even more elaborate questions can be tackled. For instance, the stories can be linked to mitigation strategies that shows how anticipated impacts might be addressed. The data analysis is therefore dependent on how the users of SSE gather their data. It could, for instance, be a qualitative, in-depth discussion on possible mitigation strategies. But it could also be accountability or importance ratings on which stakeholders should be responsible for the mitigation of impacts.

In addition, the emerging scenarios and/or risk frameworks inform further quantitative research. For example, the risk frameworks can be used as input to simulate new stories and potential policy impacts with generative AI tools (Barnett et al., 2024). We can also think of using (specific) scenarios as input/vignettes for public opinion surveys to gauge society's perspective on different future paths of AI, and thus different risks that society might experience in the future. Following this approach would answer the call to bring society into the loop (Rahwan, 2018), especially when it comes to moral choices and trade-offs between different futures and risks.

## 5 Example Scenarios

Scenarios A-D (see Figures below) show original scenarios from our empirical studies. Scenarios A and B were derived from a study focusing on the negative impact of Generative AI on the news environment (Kieslich, Diakopoulos, et al., 2024); Scenarios C and D from a study focusing on the impact – both positive and negative – from general purpose chatbots on individuals and society (Kieslich, Helberger, et al., 2024).

### Example Scenario A

Hans is a news photographer for a newspaper in Utrecht. Disasters and emergencies are part of his specialization.

In the summer of 2028, the Netherlands was ravaged by a heat wave. The forests are bone-dry and forest fires are the order of the day. On a Friday afternoon, Hans receives a report of a forest fire east of town. He gets in his car and is on location within 15 minutes. Hans shoots some spectacular photos and uploads them to the newspaper's website on the spot.

Then he sees that the website already has pictures of this fire, taken by his colleague Martina. It is not the first time she has outfoxed him. His supervisor has already reprimanded Hans once for his tardiness, threatening to fire him. How can Martina do this so quickly? Did she generate these pictures with AI?

Hans thinks: "What Martina can do, I can do too." And so he starts experimenting with using AI.

A week later, another report of a forest fire comes in. Hans, fearing he might lose his job otherwise, does not hesitate, types a few words in a text box, and clicks 'generate'. In a split second, his screen is filled with lifelike pictures of the burning forest, seen from different angles. He picks five photos, uploads them to the website, and sees that he is faster than Martina. His job seems safe.

On Monday, the police drop by the newsroom. They want to speak to Hans. One of his photos appeared to show a man wearing a blue cap. Based on this description, the police arrested someone on suspicion of arson. By an unfortunate coincidence, Hans generated an image with AI showing an innocent person, who is now detained. Hans immediately admits this, but is nevertheless fired and faces a lawsuit.

## Example Scenario B

In 2028, Robert is the only journalist working for the technology section of the New York Times. He is not a trained journalist, and was put in that role thanks to his background as a Computer Science graduate, and was openly asked to develop a system that could churn out content on a regular basis so that the rest of the team could be laid off. He used to write long form articles and content the old-fashioned way, at least at the start, only using his AI sidekick to create shorter articles that would have the purpose of enticing readers to find out more about the topics presented.

It was not long, though, before Robert realized that most people did not know the difference between what was well researched, substantiated content and the more generic and superficial output he was feeding them. He felt defeated. About one year in, he started believing that the world of information and news was moving at a frightening high speed towards being a mass of AI generated content with little human touch. He knew that the vast majority of readers were not educated enough to realize that; he knew that way too many would believe anything they would read. He knew swathes of people would not question whether an image was real. If it's on the news outlets, it must be real, right?

So it became real. Robert now reluctantly describes his role as AI Janitor. He keeps an eye on the content, sure, but has decided that he is too small of an entity to fight against the piles of repetitive drivel and deceptive images in the news. He sometimes entertains the idea of teaming up with others to start a news outlet made by real people, with no AI generated content whatsoever. He does not know if the people he might get in touch with are going to be real. He does not know if the sources they might use are real. He dreams of being a local reporter working with actual human beings in an actual physical space, only covering events that have been witnessed by real, actual people. He does not know if that is an option anymore, against the relentless and constant barrage of dubious information. He puts his feet on the desk of his home office while his AI editor generates more content, and wonders if it's not too late to buy a farm in Ohio.

## Example Scenario C

In the year 2029, AI personal chatbots had become more advanced. January 4th, 2029, also happens to be the death anniversary of the Mirandas father passing. After Mirandas father passed, she had isolated herself. She had no one to talk too and felt uncomfortable being around people.

One day, while browsing the internet, Miranda comes across an article of AI and how it has helped people immensely, not just with general advice, but in which it fills a void where it keeps one from going insane. Although Miranda was uncomfortable out in public since the death of her dad, he goal was to start going out in society again. Therefore, she started this by using AI.

Miranda researched different AI's. She wasn't sure which one to use since there were so many. So, she started to use ChatGPT. ChatGPT gave great, general advice, and one of the advice being was to pick up a hobby. "But which hobby?" Miranda says. She begins to ponder on what hobby she should pick up. ChatGPT suggests writing, which in fact used to be her favorite hobby before her father passed.

So, Miranda pursued writing again. And with the help of her AI companion, it was able to check up on her to check her progress on writing. Her AI friend gave her the push and reminder to keep going. It was a slow start, Miranda wasn't sure how to write since she lost that passion, but it was slowly coming back. Miranda started to write a story evolved around her AI that was a kids book about her computer friend. She wanted to write about this experience because she felt that AI helped her so much and that it can help other people too, despite the mysterious aura around the world of AI.

In Miranda's adventure of writing, her AI helped her immensely with editing and overall sentence structure, along with flow of her story. She was able to get a lot of her writing done in a short period of time, thanks to her AI. By the end of her novel, she had realized she felt more healed and confident to go back out in society again and meet real people. She felt nervous at first, but she felt that her AI companion had helped her with general interaction and motivation on a daily basis that she knew she could get back out there.

### **Example Scenario D**

Our story follows a young girl named Ella, set 5 years in the future. Ella, now 18 years old, has grown up in a remote rural area, in a community that limits access to some technological advances that it deems harmful. For example, Ella grew up familiar with how to use a physical library with books, rather than looking up information on the internet as her primary mode of research. Ella's community teaches the power of problem solving, self sufficiency and unique thought. She is taught to work through problems and come up with her own solutions. Ella has a sense of pride in these life skills she has learned, and she is confident that they will help her succeed in her goal of becoming a top creative advertisement professional in the future.

Ella is now on her way to college. She has chosen to attend a university in a larger, more urban area, one that has a great program for the advertisement profession she is seeking.

Ella embarks on this journey to college with great excitement. Surely the community and students in this urban area is going to be vibrant, inspiring and fresh with innovative ideas.

But unfortunately for Ella, her college experience was not quite what she was expecting. Ella executed every assignment from a place of unique thought and personal problem solving. She worked with a particular emphasis on innovative ideas and approaches, striving to offer solutions that had never been done before. To her dismay, her professors often docked her points, deeming her solutions, "not practical" or "entirely unfeasible." It seemed they lacked the vision to see what she saw.

Frustrated, Ella started inquiring to her high performing classmates about their solutions, trying to understand what high grade submissions looked like. She was again, disappointed when she examined these submissions, only to find they were mostly recycled, overdone, unoriginal ideas. Worse yet, when she spoke to a particular classmate, Aaron, she learned he and some of his friends didn't even use their own brains at all to complete the assignments. Instead, they used something she had never heard of, called AI, to complete the tasks for them. She learned all you had to do was input a prompt into a computer and it basically wrote the assignment for you, you didn't have to poses a single thought of your own.

Ella went through her college career getting used to this as the norm. She continued to submit her own unique thoughts and get average grades while her classmates used AI to submit very average and unoriginal ideas for higher praise. She had just accepted that this is the way the larger world worked but on her own merit, never stopped using her own thought process for generating ideas.

Ella graduated and went on to work for an advertising firm in an even larger city. Upon entry into the work force she was surrounded by a dozen graduates just like her who were all fighting to have their ideas noticed and excel at the firm. One particular day the management called an "all hands on deck" emergency meeting at the firm. It seemed one of their largest clients wanted some really original and cutting edge ideas for their upcoming holiday campaigns and they were not happy with anything the firm leaders had come up with. So the entire firm was tasked with coming up with something that had never been done before.

Ella's peers suggested recycled ideas over and over again, none of which made the client happy. But Ella, having never stopped exercising that muscle that let her create, was able to fill an entire white board with innovative ideas to take to the client. Ella was promoted over her peers shortly after completing that assignment and saving that client account for her firm.

The moral of the story here, is that technology can be great for some things, however it is when we rely on technology to do the work that our minds should be doing, that we fail as human kind. We should not be allowed to use technology as a crutch or a replacement for actual human thought. We need to continue to teach new generations creative problem solving as a skill and not just rely on tech to do it for us. If we continue down that road, we do nothing but regress as a society and we are doomed

## References

- Amer**, M., Daim, T. U., & Jetter, A. (2013). A review of scenario planning. *Futures: The Journal of Policy, Planning and Futures Studies*, 46, 23–40. <https://doi.org/10.1016/j.futures.2012.10.003>
- Andersen**, P. D., Hansen, M., & Selin, C. (2021). Stakeholder inclusion in scenario planning—A review of European projects. *Technological Forecasting and Social Change*, 169, 120802. <https://doi.org/10.1016/j.techfore.2021.120802>
- Barnett**, J., Kieslich, K., & Diakopoulos, N. (2024). Simulating Policy Impacts: Developing a Generative Scenario Writing Method to Evaluate the Perceived Effects of Regulation. *Proceedings of the AAAI/ACM Conference on AI, Ethics, and Society*, 7, 82–93.
- Colón Vargas**, N. (2024). Exploiting the margin: How capitalism fuels AI at the expense of minoritized groups. *AI and Ethics*. <https://doi.org/10.1007/s43681-024-00502-w>
- Diakopoulos**, N., & Johnson, D. (2021). Anticipating and addressing the ethical implications of deepfakes in the context of elections. *New Media & Society*, 23(7), 2072–2098. <https://doi.org/10.1177/1461444820925811>
- Ebers**, M. (2024). Truly Risk-Based Regulation of Artificial Intelligence—How to Implement the EU’s AI Act. Available at SSRN.
- Gellert**, R. (2020). Fundamental notions. Risk and Regulation. In *The Risk-Based Approach to Data Protection* (pp. 26–53). Oxford University Press.
- Glaser**, B., & Strauss, A. (2017). *Discovery of grounded theory: Strategies for qualitative research*. Routledge.
- Guston**, D. H. (2013). Understanding ‘anticipatory governance.’ *Social Studies of Science*, 44(2), 218–242. <https://doi.org/10.1177/0306312713508669>
- Helberger**, N. (2024). FutureNewsCorp, or how the AI Act changed the future of news. *Computer Law & Security Review*, 52, 105915. <https://doi.org/10.1016/j.clsr.2023.105915>
- Kieslich**, K., Diakopoulos, N., & Helberger, N. (2024). Anticipating impacts: Using large-scale scenario-writing to explore diverse implications of generative AI in the news environment. *AI and Ethics*. <https://doi.org/10.1007/s43681-024-00497-4>
- Kieslich**, K., Helberger, N., & Diakopoulos, N. (2024). My Future with My Chatbot: A Scenario-Driven, User-Centric Approach to Anticipating AI Impacts. *The 2024 ACM Conference on Fairness, Accountability, and Transparency*, 2071–2085. <https://doi.org/10.1145/3630106.3659026>
- Lofland**, J., Snow, D., Anderson, L., & Lofland, L. H. (2022). *Analyzing social settings: A guide to qualitative observation and analysis*. Waveland Press.
- Rahwan**, I. (2018). Society-in-the-loop: Programming the algorithmic social contract. *Ethics and Information Technology*, 20(1), 5–14. <https://doi.org/10.1007/s10676-017-9430-8>
- Ramirez**, R., & Wilkinson, A. (2016). *Strategic reframing: The Oxford scenario planning approach*. Oxford University Press.
- Tian**, E., & Cui, A. (2023). GPTZero: Towards detection of AI-generated text using zero-shot and supervised methods. *GPTZero*. <https://gptzero.me>
- Veselovsky**, V., Ribeiro, M. H., Cozzolino, P., Gordon, A., Rothschild, D., & West, R. (2023). Prevalence and prevention of large language model use in crowd work (arXiv:2310.15683). *arXiv*. <http://arxiv.org/abs/2310.15683>
- Veselovsky**, V., Ribeiro, M. H., & West, R. (2023). Artificial Artificial Artificial Intelligence: Crowd Workers Widely Use Large Language Models for Text Production Tasks (arXiv:2306.07899). *arXiv*. <https://doi.org/10.48550/arXiv.2306.07899>

## Funding

The funding for this research was provided by UL Research Institutes through the Center for Advancing Safety of Machine Intelligence

## Design & Layout

Leyla Dosch and Sara Spaargaren

## Reference suggestion

Kieslich, K., Helberger, N., & Diakopoulos, N. (2025, January 29). Scenario-based Sociotechnical Envisioning (SSE): The Guidebook. [https://doi.org/10.31235/osf.io/j5ske\\_v1](https://doi.org/10.31235/osf.io/j5ske_v1)