

THE BUSINESS'S COMMERCIAL SUCCESS SUBSTANTIALLY DEPENDS UPON:

RED FLAG NO.



RED FLAGS IN THE VALUE PROPOSITION

Using or providing algorithmic systems that make predictions, recommendations or decisions that can affect people's access to rights and resources.

- Such as:**
- Banks, insurance firms and mortgage companies using automated decision-making that results in individuals being declined credit based on their age or race
 - Recruitment companies using algorithmic systems to help employers make decisions about candidates in ways that result in certain groups such as women and ethnic minorities being disproportionality removed from recruitment processes
 - Companies offering algorithmic tools and solutions to public institutions that make decisions about:
 - Law enforcement (e.g., for use in predicting crimes based on ethnicity)
 - Immigration (e.g., for use in group profiling and prediction used in detention, deportation or monitoring decision-making)
 - [Social scoring systems](#) (e.g., for use in determining access to services, benefits or opportunities)
 - Social media companies selling profiling and targeting services that enable political campaigners to spread misinformation about opponents or election dates in ways that undermine the ability of individuals to participate in political processes without interference



HIGHER-RISK SECTORS:

- Companies that offer or make use of targeted advertising such as social media, search, websites, blogs, as well as advertisers and their media agencies
- Consumer finance and credit
- Health insurance and healthcare
- Retailers when targeting customers with product promotions
- Recruitment industry including online portals and software providers
- IT and technology companies selling algorithmic solutions to government agencies including healthcare and the criminal justice system
- Political consulting firms
- Data brokers selling data analytics as a service

THE BUSINESS'S COMMERCIAL SUCCESS SUBSTANTIALLY DEPENDS UPON:

Using or providing algorithmic systems that make predictions, recommendations or decisions that can affect people's access to rights and resources.

RED FLAG NO.



5

RED FLAGS IN THE VALUE PROPOSITION



KEY QUESTIONS FOR LEADERS TO ASK OR BE ASKED:

- Do we have evidence that algorithmic profiling delivers notably greater benefits than more traditional tools for decision-making? Have we done an expert review of how different groups may be impacted differently, if it may lead to us excluding certain groups?
- How do we assess, beyond our own existing internal knowledge, whether people's rights may be impacted as a result of algorithmic solutions we provide or deploy?
- Do we have in place the necessary technical know-how and oversight to:
 - Design, build and deploy algorithmic tools in ways that minimize discriminatory and other risks?
 - Provide sufficient transparency and explainability as to how AI algorithmic decisions were made?
 - Responsibly evaluate, procure and use algorithmic tools.
- Do we have in place the necessary evaluation processes and timelines that allow us to evaluate use of our tools by third parties to ensure we are not involved with rights impacting activities?
- If challenged, are we prepared to explain the decisions we make using these tools? Can we evidence that we are not negatively and disproportionately impacting people's right to non-discrimination, privacy and other rights?

THE BUSINESS'S COMMERCIAL SUCCESS SUBSTANTIALLY DEPENDS UPON:

Using or providing algorithmic systems that make predictions, recommendations or decisions that can affect people's access to rights and resources.

RED FLAG NO.



RED FLAGS IN THE VALUE PROPOSITION



RISKS TO PEOPLE

Key risks to people resulting from this business model include impacts on:

- A. The right to non-discrimination and associated impacts on economic, social and cultural rights
 - B. Civil and Political Rights
 - C. The Right to Effective Remedy
 - D. Privacy
- These impacts on people can arise where companies design/ deploy systems that are:
 - Trained on historical or proxy data that reflects or perpetuates social inequities;
 - Deployed at scale or rapidly with limited transparency or contestability; and/or
 - Used in contexts involving power asymmetries between decision-makers and affected individuals, in particular those in situations of vulnerability.

A. Impacts on the right to non-discrimination and associated impacts on economic, social and cultural rights, such as housing, employment opportunities, livelihoods and healthcare. The use of algorithms to automate decision-making in industries as diverse as online advertising, recruitment, healthcare, retail, and consumer finance is rarely – if ever – intended to undermine individuals' right to non-discrimination. In fact, these tools have the potential to reduce or remove human bias from decision-making. Nonetheless, the opposite can also be true. Examples include:

Social media, search and websites selling targeted advertising where:

- **Landlords have been enabled to exclude users based on race, age or gender.** This has occurred when tools allowed agents to explicitly exclude certain groups from seeing housing ads. It can also happen in more subtle ways when companies allow targeting based on categories, such as age, marital status, and ZIP code, that are de facto proxies for certain groups. A series of court cases in the United States and action by the US Department of Housing and Development, have led to many companies, including [platforms such as Facebook](#), as well as [Google](#), committing to change their policies.
- **Ads for jobs placed on search platforms reinforce stereotypes, resulting in higher paying jobs being shown or opportunities only being made available to certain populations.** For example, [a 2025 case](#) where a French regulator declared that Facebook's algorithm for placing job adverts were sexist, after an investigation found that adverts for mechanic roles skewed towards men while those for preschool teachers were targeted at women. Another example is the [2024 class-action lawsuit against WorkDay](#), an AI solution company, which was accused of disproportionately disqualifying the applications of African-Americans, individuals over the age of 40, and individuals with disabilities.

THE BUSINESS'S COMMERCIAL SUCCESS SUBSTANTIALLY DEPENDS UPON:

Using or providing algorithmic systems that make predictions, recommendations or decisions that can affect people's access to rights and resources.

RED FLAG NO.



5

RED FLAGS IN THE VALUE PROPOSITION



RISKS TO PEOPLE

- **Ads for jobs, housing or credit on search platforms being selectively shown:** A [2021 investigation](#) found that Google's advertising system allowed employers and housing advertisers to exclude people listed as nonbinary or "unknown gender" from seeing certain job, housing, or credit ads. Because the exclusion was built into the platform's settings, this meant the algorithm could systematically deny visibility of opportunities to people based on gender identity, even if the advertiser didn't intend discrimination, ultimately contributing to unequal access to work, housing, and services for non-binary people seeking work.
- **Elderly populations have been targeted with fraudulent products and services, or AI supported scams to trick them out of cash or savings,** ranging from anti-ageing products to funeral insurance and reverse mortgages. In one case, [retired, politically conservative individuals in the United States were tricked into using much of their retirement savings](#) to buy marked up gold and silver coins to "protect their money from the deep state." Even though this broke the company's rules, Facebook showed ads supporting this scheme more than 45 million times over a 21-month period. Journalists and researchers have [demonstrated](#) how AI chatbots – specifically Grok in this case - can be used to craft targeted phishing campaigns, including generating convincing scam emails or social engineering content, that successfully penetrate older adults' defenses.

Discrimination in credit and insurance decision-making, for example:

- **Where loan providers rely on algorithms to analyze credit worthiness, and such algorithms are built on existing biases in determining risk of individuals and districts.** While the algorithms are not designed to pick up race as a specific data point in the determination, such information is implicitly included due to historic biases and can result in "[automatic loan denials for individuals from marginalized communities, reinforcing racial or gender disparities](#)".
- **Where insurance companies use algorithms to set the price of cover.** In 2022, [reports alleged](#) that UK car insurance firms were using algorithms that quoted higher premiums in areas with large proportions of Black or South Asian populations. In response to similar reports in the U.S., the New York Department of Financial Services adopted [guidance](#) in 2024 for insurers on using external consumer data and artificial-intelligence systems in underwriting/pricing, to prevent discrimination.
- **Where flawed data aggravates inequalities in credit scores.** A [study](#) from 2021 shows how credit-score models are less accurate for minority and low-income borrowers, and how the underlying data quality drives part of the inequity. The study shows that the predictive tools increasingly used by companies are between 5% and 10% less accurate for these groups than for higher-income and non-minority groups, due to less accurate and available data for those groups.

THE BUSINESS'S COMMERCIAL SUCCESS SUBSTANTIALLY DEPENDS UPON:

Using or providing algorithmic systems that make predictions, recommendations or decisions that can affect people's access to rights and resources.

RED FLAG NO.



RED FLAGS IN THE VALUE PROPOSITION



RISKS TO PEOPLE

- **Recruitment industry and human resources tools that discriminate.** The recruitment industry increasingly integrates automated decision-making as part of its value proposition to employers. In this context, discrimination can occur in a range of ways that have been [well highlighted by researchers](#). [Authorities in the U.S. have also warned against discriminatory AI algorithms](#), issuing guidance to protect vulnerable people. High profile examples have included companies offering tools that:
 - **Examine social media timelines and online postings about candidates** with the risk that data which should not legally or ethically exclude an individual from a job – such as political opinion, sexual orientation or having family members convicted of a crime – ends up doing so.
 - **Use Natural Language Processing to screen out candidate resumes** that don't fit an employer's prior hiring patterns, which can perpetuate, racial, gender and other discrimination.
 - **Allow employers conducting video interviews to grade verbal responses, tone, and facial expressions against high-performing employees** potentially reinforcing biases and being unable to interpret non-white faces.
 - **Use of AI algorithms to inform promotions, salary reviews and layoffs**, [where employers rely on AI to flag employees for termination or performance issues](#), raising issues of transparency, oversight and worker rights when algorithmic systems feed into termination decisions.
- **Discrimination in healthcare.** [Healthcare professionals are increasingly looking to leverage the power of artificial intelligence](#) to achieve breakthroughs in disease detection, diagnosis and

patient care plans. Examples of risk associated with the use of AI in healthcare include the following:

- [In one case](#), an algorithmic tool sold to hospitals and insurers to predict health care needs was found to underestimate the needs of Black patients.
- [Another study](#) showed bias in algorithms predicting pain from x-rays. In this case, researchers were training the models based on physician reports of pain, and since doctors are less likely to believe marginalized people when they report pain, this algorithm replicated this bias and consistently misdiagnosed people of color.

B. Impacts on Civil and Political Rights including the right to equality before the law, freedom from arbitrary arrest, freedom of assembly, the right to information, political participation (see also Red Flag 9). For example, the commercial design and sale of tools used in:

- **Predictive Policing:** In 2019, human rights organizations, journalist and academics reported that police departments in the [United States](#) and the [United Kingdom](#) were piloting private sector tools to predict crime as a means to allocate resources with discriminating effects based on race, sexuality and age. Further, a [number of US-based firms](#) (such as Intel, Motorola, Amazon) have been accused of knowingly selling products to Chinese military, police or private surveillance companies where the products were used for a variety of purposes impacting on people's rights, including surveillance, repression, identification or censorship.

THE BUSINESS'S COMMERCIAL SUCCESS SUBSTANTIALLY DEPENDS UPON:

Using or providing algorithmic systems that make predictions, recommendations or decisions that can affect people's access to rights and resources.

RED FLAG NO.



RED FLAGS IN THE VALUE PROPOSITION



RISKS TO PEOPLE

- **Predicting Recidivism Rates in Criminal Justice:** In 2016, A commercial tool developed by U.S company Northpointe to predict the likelihood of a criminal re-offending, was [assessed by Pro Publica](#). Findings included that among other things “Black defendants were far more likely than white defendants to be incorrectly judged to be at a higher risk of recidivism, while white defendants were more likely than Black defendants to be incorrectly flagged as low risk”. [More recent studies](#) have further proposed how to address potential discriminatory outcomes of such AI solutions.
- **Facial Recognition:** The proposition facial recognition tools – [the majority of which originate in the private sector](#) - is to enable users to identify individuals by comparing their facial characteristics against a database of images. Users – [such as law enforcement agencies, airports, border control and private security companies](#) – can then act on matches where an individual has committed an offence or that they deem to be a threat. Concerns about these tools include: the risk of false positives and unfair detention ([especially where they have proven to be less accurate on non-white and non-male faces](#)), and chilling effects on freedom of assembly. In extreme cases, concerns have been raised about arbitrary detention, misidentification, including in military targeting, and the broader impact on civilian populations' rights to privacy, freedom of movement, and due process, such as in:
 - China, where deployment of [commercially developed tools](#) have been allegedly been used in [predictive policing and facial recognition in the Xinjiang region](#), which has led to descriptions of the region as an [“open air prison”](#);
 - Gaza, where the [use of commercially-developed facial recognition and biometric surveillance tools by Israel](#) includes systems used to identify and track individuals at checkpoints and during military operations;
 - The United States, where U.S. Immigration and Customs Enforcement (ICE) [deployed a variety of mobile facial-recognition and biometric identification tools](#), supplied by private firms such as Palantir and Clearview AI, in interior enforcement operations - including a smartphone app used by ICE agents to scan faces and fingerprints against federal databases - prompting lawsuits and congressional scrutiny over issues pertaining to privacy, misidentification, and civil liberties.
- **Political Campaigning and Disinformation:** Social media companies that generate revenue by selling targeted advertising to political campaigners have come under intense scrutiny from [civil society organizations making the case that this has threatened democratic processes](#). Of particular concern have been examples in which [voters have been targeted by foreign parties with disinformation](#) about voting dates and processes with the aim of suppressing some voters from going to the polls. Equally concerning, and spotlighted by the infamous [Cambridge Analytica scandal](#), are when political lobby or consulting firms sell micro-targeting strategies using disinformation as a service to political incumbents or opposition parties.

THE BUSINESS'S COMMERCIAL SUCCESS SUBSTANTIALLY DEPENDS UPON:

Using or providing algorithmic systems that make predictions, recommendations or decisions that can affect people's access to rights and resources.

RED FLAG NO.



5

RED FLAGS IN THE VALUE PROPOSITION



RISKS TO PEOPLE

- **Lobby groups and Disinformation:** Similarly, lobby groups seeking to influence citizen behavior, including with respect to their own healthcare, can take advantage of targeted advertising. For example, [Google has been accused](#) of promoting misleading search ads from “fake” abortion clinics, funded by anti-choice lobby groups, that appear to offer reproductive healthcare but actually aim to dissuade people from accessing abortions, thereby misdirecting individuals seeking legitimate care. As it is difficult for individuals to easily distinguish ads from unbiased information, this ad targeting intercepts vulnerable individuals at critical moments, undermining their ability to exercise their reproductive rights and impedes access to accurate healthcare information.

C. Impacts on the Right to Effective Remedy: Whether algorithmic profiling and predictions amount to State violations, or a business abuse of human rights, the nature of the tools described above can undermine the right to an effective remedy for violations of human rights, which is a fundamental principle of international human rights law. In her [2020 report](#), the UN Special Rapporteur on contemporary forms of racism, racial discrimination and xenophobia explains that, “In many cases, the data, codes and systems responsible for discriminatory and related outcomes are complex and shielded from scrutiny, including by contract and intellectual property laws. In some contexts, not even computer programmers may themselves be able to explain the way that their algorithmic systems function. This “black box” effect makes it difficult for affected groups to overcome steep evidentiary burdens of proof typically required to

prove discrimination through legal proceedings, assuming that court processes are even available in the first place.”

D. Privacy Impacts: Where business models depend on selling or deploying capacities for algorithmic profiling and prediction about individuals, this can create or compound risks to the right to privacy (see also Red Flag 16). For example:

- **By revealing details about an individual's private life against their will.** In 2012, U-S retailer Target [was criticized for posting coupons for baby products to the household email account used by a teenage girl](#), thereby alerting her father to the fact of her pregnancy. The company's predictive analytics tools had deciphered that the girl was pregnant based on her recent purchase history. [A more recent paper](#) points out how outside of the healthcare system, even seemingly innocuous personal data can be used to infer, and reveal, private information, including reproductive health information. Some [individuals have reported that adverts clearly targeting members of the LGBTQ community](#) have appeared on their newsfeed in sight of family members or colleagues, resulting in the individual feeling outed.
- **By incentivizing “data maximization.”** Algorithms need to be trained on vast amounts of training data concerning the personal habits, preferences and behaviors of individuals. This may incentivize developers and engineers building algorithmic systems to collect or source data without putting in place privacy-protecting processes e.g. around consent.

THE BUSINESS'S COMMERCIAL SUCCESS SUBSTANTIALLY DEPENDS UPON:

Using or providing algorithmic systems that make predictions, recommendations or decisions that can affect people's access to rights and resources.

RED FLAG NO.



5

RED FLAGS IN THE VALUE PROPOSITION



RISKS TO THE BUSINESS

- **Rapidly Evolving Regulatory Risks:** The development, sale and use of algorithmic profiling and decision-making tools is gaining increased attention from regulators. Some examples include:
 - The EU's General Data Protection Regulation addresses the right of individuals not to be data subject to a decision based solely on automated processing, including profiling, where that decision has legal or other effects concerning him or her or similarly significantly affects them. In one example, a [Swedish financial services company](#) was ordered to correct its credit risk algorithm which was illegally using age as a parameter to determine credit.
 - The [EU AI Act](#) includes obligations for "high-risk" AI systems, including those that underpin AI-enabled business models. High-risk AI systems include systems used for recruitment and performance evaluation, credit scoring, insurance risk assessment, emotion recognition, biometric identification, and critical infrastructure applications. Violations may trigger fines of up to €35 million or 7% of annual worldwide turnover, whichever is higher.
 - Enacted in 2021, New York City's [Local Law 144](#) prohibits employers and employment agencies from using an automated employment decision tool unless the tool has been subject to a bias audit within one year of the use of the tool, information about the bias audit is publicly available, and certain notices have been provided to employees or job candidates.
- **Existing Legal Risk:** Where algorithms are being designed and used to make traditional decisions in novel ways concerning employment, advertising and credit, existing laws apply. For example:
 - The use of [AI in hiring in the United States](#) may lead to [companies failing to comply with existing laws](#) such as the Employee Polygraph Protection Act or Genetic Information Non Discrimination Act.
 - Plaintiffs in several [lawsuits in the US](#) have claimed that the use of AI-powered hiring software discriminates against older candidates and people of color, by rejecting such candidates more frequently.
 - In the UK, law firms and academic institutions have warned that financial institutions that make use of algorithms can risk non-compliance with consumer lending laws.
- **Reputational Risk:** The sharp increase in civil society scrutiny of algorithmic tools means that companies developing or using such tools may experience reduced trust from consumers, employees and citizens. For example:
 - **Palantir** has [faced protests and public criticism](#) over its work with ICE, including demonstrations outside its offices in New York City organized by immigrant rights groups citing concerns about surveillance and deportations tied to its tools. [Additionally, AI generated misinformation can quickly spread on the internet, damaging consumer trust and ultimately affecting a company's bottom line.](#)

THE BUSINESS'S COMMERCIAL SUCCESS SUBSTANTIALLY DEPENDS UPON:

Using or providing algorithmic systems that make predictions, recommendations or decisions that can affect people's access to rights and resources.

RED FLAG NO.



5

RED FLAGS IN THE VALUE PROPOSITION



RISKS TO THE BUSINESS

- [In 2025](#), the Polish government called on the EU Commission to investigate **TikTok** after the social media platform hosted artificial intelligence-generated content including calls for Poland to withdraw from the EU, which the government claimed was a Russian disinformation campaign.
- **Lost Investment Pre-Launch:** Where algorithmic systems are found to discriminate or are deemed to be making decisions in ways that lack social license, this can result in companies having to choose not to take these products to market.
 - In 2018, one company had to [halt the launch of a product](#) designed to vet people for domestic services using “advanced artificial intelligence” to analyze their personalities based on social media posts, after they faced a public backlash.
 - A lucrative [commercial partnership between Amazon Ring and surveillance firm Flock Safety](#) was abandoned following widespread public and regulatory backlash ignited by a Super Bowl advertisement, in which the partnership was perceived to be promoting invasive surveillance capabilities. In a letter to the Amazon CEO, a US Senator stated that in airing this commercial, Amazon had “inadvertently revealed the serious privacy and civil liberties risks attendant to these types of Artificial Intelligence-enabled image recognition technologies.”

Using or providing algorithmic systems that make predictions, recommendations or decisions that can affect people's access to rights and resources.



5

RED FLAGS IN THE VALUE PROPOSITION



WHAT THE UN GUIDING PRINCIPLES SAY:

**For an explanation of how companies can be involved in human rights impacts, and their related responsibilities, see [here](#).*

Companies that make decisions and pursue actions based on algorithmic profiling and predictions can **cause** adverse impacts on human rights. An example would be a bank denying credit based on a tool that makes discriminatory recommendations, or the use of algorithms to process user data in ways that result in privacy infringements.

Companies whose value proposition is to sell the capability to profile and predict to public or private third parties can **contribute to** adverse human rights impacts that those actors cause where their tools embed discriminatory biases or are designed in ways that enable or incentivise other adverse human rights impacts. Contribution might arise due to the ways that customers are empowered to use these tools (such as by excluding certain groups) or may be more subtle such as when an algorithmic system has bias built into the data set.

An added complexity is that a single algorithmic system may integrate a number of inputs from different actors. For example, a data broker might provide training data; an AI research firm might license an algorithm and a developer might design the customer interface. Depending on the specific circumstances, each of these companies could contribute to adverse impacts.

In situations where companies have taken reasonable steps to prevent their tools contributing to discrimination and other human rights harms, they may nevertheless be **linked to** adverse impacts that business or government customers are causing.



POSSIBLE CONTRIBUTIONS TO THE SDGs:

Algorithmic systems may be used to advance **a number of SDGs** such as those listed below. Addressing impacts to people associated with this red flag can contribute to ensuring that this is done in ways that do not simultaneously impact people's rights to non-discrimination, privacy and physical and mental health and well-being.



SDG10: Reduce Inequality within and Among Countries.



SDG3: Healthy Lives and Well-Being for all, including by tackling disruptions to progress such as from the COVID-19 global pandemic.



SDG 5.B: Promote Empowerment of Women Through Technology



SDG11: Make cities and human settlements inclusive, safe, resilient

The [UN Secretary-General's Roadmap for Digital Cooperation](#) is an important resource to guide "all stakeholders to play a role in advancing a safer, more equitable digital world" even as technological solutions are used to achieve the SDGs.

THE BUSINESS'S COMMERCIAL SUCCESS SUBSTANTIALLY DEPENDS UPON:

Using or providing algorithmic systems that make predictions, recommendations or decisions that can affect people's access to rights and resources.

RED FLAG NO.



5

RED FLAGS IN THE VALUE PROPOSITION

DUE DILIGENCE LINES OF INQUIRY:

Unless otherwise indicated, the following questions draw heavily on [Ranking Digital Rights' Best Practices: Algorithms, Machine Learning and Automated Decision-Making](#), and the World Economic Forum's [White Paper on How to Prevent Discriminatory Outcomes in Machine Learning](#).

The following questions can guide due diligence on the part of a company using algorithmic decision making, as well as provide a guide for companies selling such tools to assess the capacity and commitment of their customers or ultimate end users to respect rights in their use.

- Do we have a clear policy that describes how the company identifies and manages human rights risks related to the algorithmic system(s) we use?
- Do we inform customers or users about the existence of algorithmic profiling, describe how this works, explain the variables that influence the algorithm, and explain how users and customers may be impacted?
- Have we mapped and understood if any particular groups may be at an advantage or disadvantage in the context in which the system is being deployed? Do we have a method for checking if the output from an algorithm is decorrelated from protected or sensitive features?
- Do we seek a diversity of views about the potential risks of proposed models, especially from specific populations affected by the outcomes of algorithmic systems we use?
- Have we established robust diversity and inclusion policies at every level of the company, and notably in teams that develop algorithms, machine learning models, or other automated decision-making tools?
- Have we consulted with all the relevant domain experts whose interdisciplinary insights allow us to understand potential sources of bias or unfairness, and to design ways to counteract them?
- Do we assess whether any uses or use-cases of our algorithmic tools pose risks to human rights? Where we identify these, are we:
 - Creating clear and enforceable terms of use?
 - Engaging enterprise and government customers/users to educate and train them about how to use the tools without increasing human rights risks?
 - Do we have systems in place to monitor and review how customers are using our tools?
 - Are we clear about the actions we will take if we discover that our tools are being used in ways that lead to, or increase the likelihood of, adverse human rights impacts?
- Do we apply "rigorous pre-release trials to ensure that algorithmic systems will not amplify biases and error due to any issues with the training data, algorithms, or other elements of system design?"
- Have we outlined an ongoing system for evaluating fairness throughout the life cycle of our product? Do we have an escalation/emergency procedure to correct unforeseen cases of unfairness when we uncover them?

THE BUSINESS'S COMMERCIAL SUCCESS SUBSTANTIALLY DEPENDS UPON:

RED FLAG NO.

Using or providing algorithmic systems that make predictions, recommendations or decisions that can affect people's access to rights and resources.



5

RED FLAGS IN THE VALUE PROPOSITION

DUE DILIGENCE LINES OF INQUIRY:

- Are we clearly committed to only buying and/or using training datasets that comprise data whose data subjects have provided informed content to having their data included in datasets used for this purpose?
- Are we making dataset(s) used to train machine learning models, terms of use, and APIs available to allow third parties to provide and review the behavior of our system?
- What reporting, grievance or redress processes and recourse do we have in place? Do we have a process in place to make necessary fixes to the design of the system based on reported issues or concerns?
- What measures can we build into our product to limit risky use or limit resultant impacts?
- What are our triggers and processes for considering and effecting a responsible exit from the relationship should concerns increase?

[On risks arising from use of products and business partnerships, see Red Flags 9 and 15 respectively]

Additional questions for tool developers regarding sale to third parties:

- What existing impacts and vulnerabilities in the geography of roll-out can be exacerbated by the use of our tool/technology? (e.g. What groups face existing discrimination or other marginalization or violence and how might these impacts be exacerbated or replicated by use of our tool?)
- What is our assessment of the quality of the customer/user's human rights due diligence, especially their identification of their own salient human rights issues, ongoing stakeholder engagement and human rights-related governance? Can we influence – including require - that party to improve their processes?

HIGH-LEVEL
DECISION-MAKING

RISK TO PEOPLE

RISK TO THE
BUSINESS

UNGP's & SGD's
ANALYSIS

TAKING ACTION

Using or providing algorithmic systems that make predictions, recommendations or decisions that can affect people's access to rights and resources.



5

RED FLAGS IN THE VALUE PROPOSITION



MITIGATION EXAMPLES:

Mitigation examples are current or historical examples for reference, but do not offer insight into their relative maturity or effectiveness. Moreover, some examples listed below are proposals for mitigating actions that have come from data science and engineering research institutes.

- **Principles, Governance and Oversight:** A number of companies in the technology industry and beyond have committed to some form of AI fairness principles as well as having ethics officers and cross-functional committees that look at these issues. One example is Microsoft's AI, Ethics and Effects of Engineering and Research (AETHER) Committee, which operates alongside the company's Office of Responsible AI (ORA). [Microsoft states](#) that its governance arrangements are designed to set "company-wide rules for enacting responsible AI, as well as defining roles and responsibilities for teams involved in this effort" and that "senior leadership relies on Aether to make recommendations on responsible AI issues, technologies, processes, and best practices." Additionally, [companies can adopt end-to-end governance frameworks](#) that continuously monitor models for bias, use domain-specific training data and diverse datasets to improve fairness, and conduct red-teaming and internal audits to catch discriminatory outputs.
- **Tech Tools to Detect Bias:** Some technology companies – including [IBM](#), [Microsoft](#), [Google's What-If tool](#) and [Facebook's Fairness Flow](#) – have developed products aimed at detecting bias in algorithmic decision-making. Such efforts can be a way to root-out bias from companies' own profiling and predictive models as well as a way for "big tech" to mitigate the risk that third parties develop, design and deploy discriminatory algorithms using these companies' platforms or computing power. With a similar purpose, [Aequitas](#) is an "an open source bias audit toolkit developed at the University of Chicago, [that] can be used to audit the predictions of machine learning based risk assessment tools to understand different types of biases, and make informed decisions about developing and deploying such systems."
- **Debiasing Discrimination in Lending:** Start-up [Zest AI](#) has [created a feature](#) that "uses a technique called adversarial debiasing to correct discrimination in lending models... One model predicts a borrower's ability to pay, while the second predicts protected information, such as the race or gender of the borrower. The dueling models learn from each other through dozens of adjustments until the discrimination predictor is stumped — the race or gender variable bears no meaningful relationship to the applicant's credit score."
- **Hiring and talent assessment tools:** Some AI platforms (e.g., [Knockri](#)) use AI evaluation methods designed to reduce bias in hiring by focusing on skills and transcripts, not demographic cues like appearance or voice.
- **Data-sheets for Data Sets:** Studies have proposed the idea of [labelling of data sets that train algorithms similar to a nutrition labelling on foods](#). The intent would be to mitigate against discriminatory outcomes that occur when biased data sets are used to train algorithmic models. The idea is that this will "allow users to understand the strengths and limitations of the data that they're using and guard against issues such as bias and overfitting."

THE BUSINESS'S COMMERCIAL SUCCESS SUBSTANTIALLY DEPENDS UPON:

RED FLAG NO.

Using or providing algorithmic systems that make predictions, recommendations or decisions that can affect people's access to rights and resources.



RED FLAGS IN THE VALUE PROPOSITION



MITIGATION EXAMPLES:

- **Changes to Targeted Advertising Policies:** As far back as 2015, [Facebook and Google banned payday loan companies from advertising](#) on their platforms. Since 2019, TwitterX, Google and Facebook have made changes to the policies and systems that allow customers to target adverts. Different changes pertain to different categories of advert including for housing and job opportunities. Of particular interest from a human rights perspective were the changes that Twitter and Google made to policies concerning political advertising made in the run-up to the 2020 US presidential election. [Twitter banned political ads](#) outright in October 2019, while [Meta decided in 2025 to ban all political, electoral and social issue ads](#), in response to EU regulation, which is also the case for [Google](#).
- **LinkedIn Fairness Tool Kit:** [Linked-In has developed LiFT](#) an open-source project that detects, measures, and mitigates biases in training data sets and algorithms. The company has been using the tool itself to “compute the fairness metrics of training datasets on its platforms, such as the Job Search model.”
- **Ideal's Reduce-Bias Guidance and tool to Reduce Bias:** The recruitment services firm Ideal published a [Workplace Diversity Through Recruitment: A Step-By-Step Guide](#) and has a tool that customers can use to test and monitor for adverse impacts in its candidate grading system. Customers who collect demographic data during the course of their hiring

process, can ask Ideal to instruct its algorithms to both ignore those demographics and test for and remove adverse impacts based on, among other things, compliance with the US Department of Labor's affirmative action program, Canada's equity programs for designated groups, and the European Union's hiring discrimination laws.

HIGH-LEVEL
DECISION-MAKING

RISK TO PEOPLE

RISK TO THE
BUSINESS

UNGP's & SGD's
ANALYSIS

TAKING ACTION

THE BUSINESS'S COMMERCIAL SUCCESS SUBSTANTIALLY DEPENDS UPON:

Using or providing algorithmic systems that make predictions, recommendations or decisions that can affect people's access to rights and resources.

RED FLAG NO.



5

RED FLAGS IN THE VALUE PROPOSITION



OTHER TOOLS AND RESOURCES:

- Ranking Digital Rights' Best Practices (2019): [Algorithms, Machine Learning and Automated Decision-Making](#)
- World Economic Forum (2018), [White Paper on How to Prevent Discriminatory Outcomes in Machine Learning](#).
- OECD (2025), [Governing with Artificial Intelligence: The State of Play and Way Forward in Core Government Functions](#)
- European Data Protection Supervisor (2025), [TechDispatch #2/2025 – Human Oversight of Automated Decision-Making](#)
- Council of Europe (2023), [Algorithmic Transparency and Accountability of Digital Services](#)
- [Racial Discrimination and Emerging Digital Technologies: A Human Rights Analysis](#) (A/HRC/44/57), Report from the UN Special Rapporteur on contemporary forms of racism (2020)
- European Union Agency for Fundamental Rights (2018) [Big Data: Discrimination in data-supported decision making](#).
- Harvard School of Law and Corporate Governance (2020) [Artificial Intelligence and Ethics: An Emerging Area of Board Oversight Responsibility](#).
- UK Information Commissioners Office [GDPR: Rights Related to Automated Decision-Making including Profiling](#).
- [UN Secretary-General's Roadmap for Digital Cooperation](#).

This resource is part of Shift's collection of Business Model Red Flags, developed as part of the Valuing Respect Project and generously funded by Ministry of Foreign Affairs Finland, the Norwegian Ministry of Foreign Affairs, and Norges Bank Investment Management. The resource was updated during 2025-26 with generous funding from Generation Foundation. Learn more at: shiftproject.org/valuing-respect