

STRATEGIC FILTERING FOR CONTENT MODERATION: FREE SPEECH OR FREE OF DISTORTION?

Anonymous authors

Paper under double-blind review

ABSTRACT

User-generated content (UGC) on social media platforms is vulnerable to incitements and manipulations, necessitating effective regulations. To address these challenges, those platforms often deploy automated content moderators tasked with evaluating the harmfulness of UGC and filtering out content that violates established guidelines. However, such moderation inevitably gives rise to strategic responses from users, who strive to express themselves within the confines of guidelines. Such phenomena call for a careful balance between: 1. ensuring freedom of speech — by minimizing the restriction of expression; and 2. reducing social distortion — measured by the total amount of content manipulation. We tackle the problem of optimizing this balance through the lens of mechanism design, aiming at optimizing the trade-off between minimizing social distortion and maximizing free speech. Although determining the optimal trade-off is NP-hard, we propose practical methods to approximate the optimal solution. Additionally, we provide generalization guarantees determining the amount of finite offline data required to approximate the optimal moderator effectively.

1 INTRODUCTION

The Internet has fostered a global ecosystem of social interaction, where social media plays a central role. People use platforms to connect with others, engage with news content, share information, and entertain themselves. Yet, the nature of online content and interactions has increasingly raised concerns among policymakers. Social media can be weaponized to advance extremist causes or disseminate misinformation and fake news. For example, social bots were used to manipulate public discourse during the 2016 U.S. presidential election (Bessi & Ferrara, 2016; Boichak et al., 2018). Another growing concern is cyberbullying, which disproportionately targets public figures and targeted groups (O’Regan & Theil, 2020; West, 2014; MICHAEL, 2024). In response to these issues, several European governments have introduced regulations aimed at curbing fake news and hate speech on social platforms (Cauffman & Goanta, 2021). These incidents highlight the need for robust content moderation mechanisms—policies or algorithmic tools implemented by platforms to limit abuse, promote healthy discourse, and enhance user experience (Roberts, 2018; Barocas & Selbst, 2016). However, while content moderation is crucial for managing platform integrity, it also raises significant concerns regarding censorship and freedom of expression (Shaughnessy et al., 2024; Bollinger & Stone, 2022; Brannon, 2019; Young, 2022). Overly restrictive moderation may inadvertently suppress legitimate content and discourage diversity of expression. In this work, we investigate how to algorithmically balance the competing goals of mitigating harmful content and preserving free speech.

A further complication arises from users’ strategic responses to public moderation policies. It is widely observed that users actively engage with *harmful social trends*—not only to gain visibility, but also to evade platform enforcement by subtly modifying their content. For instance, users often manipulate hashtags or obscure key terms to disguise the true nature of their posts (Gerrard, 2018). Inspired by recent work in

strategic machine learning (Hardt et al., 2016), we model the interaction between platform moderators and users as a Stackelberg game, explicitly characterizing users’ strategic behavior in response to moderation when engaging with harmful trends. The platform’s objective is to design content moderation mechanisms that minimize engagement with such trends while accounting for potential strategic manipulations, and to do so in a way that avoids unnecessary content removal and respects individual rights to free expression.

Our work is related to the growing line of research on strategic ML that studies learning from data provided by strategic agents (Dalvi et al., 2004; Dekel et al., 2008; Brückner & Scheffer, 2011). Hardt et al. (2016) introduced the problem of *strategic classification* as a game between a mechanism designer that deploys a classifier and an agent that best responds to the classifier by modifying their features at a cost. Follow-up work studied different variations of this model, in a PAC-learning setting (Sundaram et al., 2023), online learning (Dong et al., 2018; Chen et al., 2020; Ahmadi et al., 2021), incentivizing agents to take improvement actions rather than gaming actions (Kleinberg & Raghavan, 2020; Haghtalab et al., 2020; Alon et al., 2020; Ahmadi et al., 2022), causal learning (Bechavod et al., 2021; Perdomo et al., 2020), fairness (Hu et al., 2019), etc.

1.1 OUR RESULTS AND CONTRIBUTIONS:

User agency modeling. We propose a behavioral model that captures the agency of social media users engaging with social trends to gain visibility while staying within moderation boundaries, enabling platform designers to frame the optimal moderation as a mechanism design problem.

Optimizing the trade-off. We reduce the mechanism design problem of balancing the goals of curbing harmful trends and protecting free speech to a constrained optimization formulation with offline data. We then propose an effective approximation method to compute the optimal content moderator and validate its effectiveness through experiments on synthetic environments.

Statistical and computational hardness results. We provide two key theoretical results: (1) a sample complexity bound (Theorem 1) showing that solving the problem using a polynomial number of offline samples leads to good generalization performance across a broad class of moderation functions (Proposition 3); and (2) a computational hardness result establishing that even for the simple linear function class, finding the optimal content moderator is NP-hard (Theorem 2). This inherent computational complexity motivates our approximation-based approach.

Conceptual contribution. Beyond technical results, our work offers a key insight for welfare-driven platform designers: the best practice for mitigating harmful social trends is not to simply suppress associated content, but to define a moderation criterion with a carefully chosen threshold such that a large fraction of content lies near its decision boundary. We formalize this intuition, fully characterize the computational challenges, and propose practical approximation methods to address them.

2 PROBLEM SETTING

Let $\mathcal{X} \subset \mathbb{R}^d$ denote the feature space of each user’s generated content (UGC). Throughout the paper we assume that $\mathcal{X}$ is convex and compact. Our problem formulation is built upon the interplay between a set of n users on a social media platform and an automated content moderator $\mathcal{M}$, as shown in the following.

User representation: A user indexed by i is represented by a tuple $\mathbf{u}_i = (\mathbf{x}_i, c_i)$, where $\mathbf{x}_i \in \mathcal{X}$ is the feature vector of $\mathbf{u}_i$ ’s generated content representing her original intention of expression and c_i denotes the manipulation cost. We consider the case where the user wants to tweak the original message $\mathbf{x}_i$ to $\mathbf{z}_i$ to better align with a global ongoing social trend $\mathbf{e}$, at a marginal cost c_i .

Content moderator: The role of a content moderator $\mathcal{M}$ is to regulate published content, ensuring it adheres to platform guidelines. Without loss of generality, $\mathcal{M}$ can be regarded as an indicator function $\mathbb{I}[f(\mathbf{x}; \mathbf{w}) \leq 0]$, where 0 indicates that the content is flagged as problematic and should be moderated, while 1 indicates it is benign. For simplicity, we define the content moderator as the function $f(\mathbf{x}; \mathbf{w}) : \mathbb{R}^d \rightarrow \mathbb{R}$, parameterized by $\mathbf{w}$, and refer to the set $\mathbf{x} \in \mathcal{X} : f(\mathbf{x}; \mathbf{w}) \leq 0$ as the benign region associated with f . The output of f can be interpreted as a harmfulness score for each content. In this work, we will focus on moderators f that induce convex benign regions.

User’s strategic response: With the components outlined above, we can formulate a utility function to capture the potential strategic behavior of a user and predict her response when facing a moderator $f(\cdot; \mathbf{w})$, given her profile $\mathbf{u} = (\mathbf{x}, c)$.¹ The following Eq. (1) characterizes the user’s utility when modifying her published content from $\mathbf{x}$ to $\mathbf{z}$:

$$u(\mathbf{z}; (\mathbf{x}, c), \mathbf{e}, f) = \mathbb{I}[f(\mathbf{z}) \leq 0] \cdot \mathbf{z}^\top \mathbf{e} - c \|\mathbf{z} - \mathbf{x}\|^2, \quad (1)$$

where the first term reflects the gain by aligning with trend $\mathbf{e}$ while not being moderated by f , and the second terms measures the manipulation cost. The following proposition 1 characterizes the property of the best response $\mathbf{z}^*$ that maximizes Eq. (1).

Proposition 1. Denote the best response of user $\mathbf{u} = (\mathbf{x}, c)$ against a convex content moderator $f(\mathbf{z}; \mathbf{w})$ by

$$\mathbf{z}^* = \Delta(\mathbf{x}, c; \mathbf{e}, f) = \arg \max_{\mathbf{z} \in \mathcal{X}} u(\mathbf{z}; (\mathbf{x}, c), \mathbf{e}, f), \quad (2)$$

and let $\mathbf{z}' = \mathbf{x} + \frac{\mathbf{e}}{2c}$. Then $\mathbf{z}^*$ always exists and has the following characterizations:

1. if $f(\mathbf{z}') \leq 0$, $\mathbf{z}^* = \mathbf{z}'$.
2. if $f(\mathbf{z}') > 0$ and $f(\mathbf{x}) \leq 0$, $\mathbf{z}^* = \mathcal{P}_f(\mathbf{z}')$, where $\mathcal{P}_f(\mathbf{x})$ denotes the ℓ_2 projection of $\mathbf{x}$ on to the hyperplane $\{\mathbf{x} \in \mathbb{R}^d : f(\mathbf{x}) = 0\}$.
3. if $f(\mathbf{z}') > 0$ and $f(\mathbf{x}) > 0$, $\mathbf{z}^* = \mathbf{x}$ or $\mathcal{P}_f(\mathbf{z}')$, depending on the location of $\mathbf{x}$.

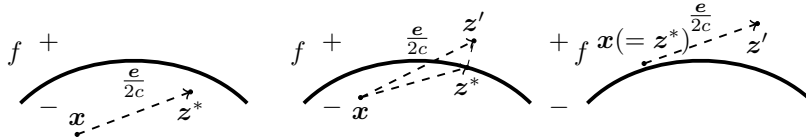


Figure 1: Illustration of best response $\mathbf{z}^*$. Left: if $\mathbf{z}' = \mathbf{x} + \frac{\mathbf{e}}{2c}$ is benign, $\mathbf{z}^* = \mathbf{z}'$. Middle: if $\mathbf{x}$ is benign but $\mathbf{z}'$ is problematic, $\mathbf{x}$ moves to the projection of $\mathbf{z}'$ on f . Right: if $\mathbf{x}$ is already problematic, $\mathbf{z}^* = \mathbf{x}$, or the projection of $\mathbf{z}'$ on f , depending on which one yields a higher utility.

Proposition 1 reveals a two-level response pattern. In the first level, each content $\mathbf{x}$ tends to shift towards an idealized location $\mathbf{z}' = \mathbf{x} + \frac{\mathbf{e}}{2c}$, manipulating its features in the trending direction $\mathbf{e}$ by an amount determined by the cost. Such $\mathbf{z}'$ is also the user’s preferred manipulation result in the absence of any moderation. The second level can be viewed as a self-correction process starting from $\mathbf{z}'$: if $\mathbf{z}'$ is accepted by the moderator f , it becomes the user’s final response; however, if $\mathbf{z}'$ is flagged as problematic by f , the user would adjust it to the closest point on f ’s decision boundary, ensuring minimal alteration while still complying with the platform’s guidelines. The proof of Proposition 1 is deferred to Appendix B. Proposition 1 highlights the

¹Our utility model connects to previous works, e.g., if we replace the term $\mathbf{z}^\top \mathbf{e}$ with a user-dependent preference parameter r , our model reduces to the agent utility proposed in (Sundaram et al., 2023).

role of content moderation in reducing distortions introduced by the trending direction e , which may deviate from users' true expressive intent, represented by x . For UGC near the moderation boundary, moderation can mitigate distortion by incentivizing users to align their content with platform guidelines. Thus, the platform can intuitively reduce overall social distortion by encouraging more UGC to move closer to this boundary. However, this strategy comes with a trade-off: the risk of filtering out certain UGC, potentially infringing on users' freedom of expression. This presents a key challenge for the platform—how to balance reducing social distortion with preserving free speech. In the following section, we formalize this problem and explore its complexity and possible solutions.

3 THE SOCIAL DISTORTION, FREEDOM OF SPEECH, AND THEIR TRADE-OFF

In this section, we formally introduce the concept of social distortion and explain how content moderation reduces social distortion but at the potential cost of infringing on freedom of speech. From Proposition 1, we observe two major effects of deploying f : first, it discourages users from excessively following social trend e , which is beneficial; second, it may flag some UGC as harmful, potentially leading to user churn. The first effect can be quantified by the displacement of users who were initially on the benign side. The second effect can be assessed by the proportion of UGC remaining on the platform, serving as an index for freedom of speech. The following Definition 1 formally introduces the concept of social distortion (SD):

Definition 1. *The social distortion (SD) of moderator f induced on a user (x, c) is defined as*

$$D(f; (x, c), e) = \begin{cases} \|x - \Delta(x, c; e, f)\|_2^2, & \text{if } f(x) \leq 0, \\ 0, & \text{otherwise.} \end{cases} \quad (3)$$

The social distortion function D defined in Eq. (3) quantifies the manipulation effort of x , which measures how much the user's strategic adaptation diverges from her true expressive intent. Importantly, our definition of social distortion applies only to UGC whose original representation x is not filtered by f . This is because the strategic behavior of users with $f(x) > 0$ does not contribute to the distortion negatively: as Proposition 1 suggests, users with $f(x) > 0$ are either filtered out—meaning their content is not distorted—or they adjust x to align with platform guidelines, which is considered a beneficial manipulation and thus should not be counted as distortion. Following Definition 1, an immediate observation is that for any user (x, c) , deploying f does not increase social distortion relative to an unmoderated environment, as substantiated by the following Proposition 2.

Proposition 2. *Let $\perp$ denote a trivial moderator that marks all $x \in \mathcal{X}$ as benign. It holds that*

$$D(f; (x, c), e) \leq D(\perp; (x, c), e), \quad (4)$$

and the inequality holds strictly if and only if $f(x) \leq 0 < f(x + \frac{e}{2c})$.

Proposition 2 illustrates a nontrivial moderator can always reduce a user's distortion in her expression, as it requires adjusting x to stay closer to the boundary of f compared to the response $x + \frac{e}{2c}$ induced by a trivial moderator. Based on this, we propose a natural optimization objective that evaluates the expected social distortion mitigated by f across a population of users:

Definition 2. *The social Distortion Mitigation (DM) induced by a moderator f over a user $u = (x, c)$ is the difference between the average social distortion induced by a trivial moderator $\perp$ and f on u , i.e.,*

$$h(f; e, x, c) = D(\perp; (x, c), e) \cdot \mathbb{I}[f(x) \leq 0] - D(f; (x, c), e), \quad (5)$$

and the total social distortion mitigation induced by f on a population of users $U = \{u = (x_i, c_i)\}_{i=1}^n$ is thus defined as

$$DM(f; U) = \sum_{u \in U} h(f, u), \quad (6)$$

where $\Delta(x, c; f, e)$ is defined in Eq. (2).

Given a class of candidate moderator functions $\mathcal{F} = \{f\}$ and a user distribution $\mathcal{U}$, the problem of optimizing expected social distortion over $\mathcal{U}$ can be formulated as finding an $f \in \mathcal{F}$ that maximizes $DM(f; \mathcal{U})$. However, there is no guarantee on how much freedom of speech the optimal moderator f will sacrifice—that is, how many users may need to be filtered out. The trivial moderator $f = \perp$, which does not filter out any users, cannot mitigate any social distortion. This suggests that any moderator aiming to reduce a reasonable amount of social distortion must inevitably sacrifice some degree of freedom of speech.

To illustrate this trade-off, consider the toy model in Figure 2, where $\mathbf{x}$ is uniform distributed in a unit ball centered at the origin, and the social trend is $\mathbf{e} = (1, 0)$. Clearly, any reasonable moderator f that maximizes DM would have a decision boundary perpendicular to $\mathbf{e}$, as this direction maximizes the deterrent effect of f on users’ strategic responses. For each moderator f of the form $x = \theta, \theta \in [-1, 1]$, we can plot the induced social distortion mitigation and a freedom of speech preservation index, which is the fraction of content still allowed on the platform, as shown in the right panel of Figure 2. As f moves from the left margin of $\mathcal{X}$ to the right margin, the social distortion mitigation exhibits an inverted U-shape, while the freedom of speech index consistently increases. This illustrates the trade-off between the two measures. The tension arises because the maximum social distortion mitigation is intuitively achieved when f is positioned where the content distribution is most concentrated, whereas freedom of speech preservation pushes the optimal f toward the margins of the distribution, making it difficult to achieve a doubly optimal moderator. As we

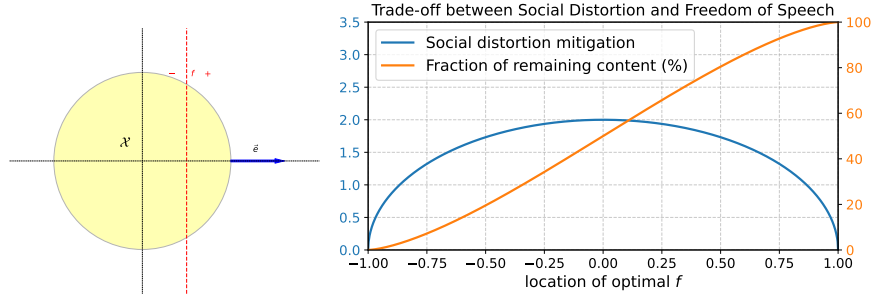


Figure 2: An illustration of the trade-off between social distortion and freedom of speech in a toy model. Left: The original UGC distribution is uniformly random within a unit ball in $\mathbb{R}^2$, with a social trend $\mathbf{e} = (1, 0)$. Right: The optimal function f under varying freedom of speech constraints and the resulting induced social distortion.

can learn from the toy example, the key challenge is determining how to strike a balance between the social distortion and freedom of speech objectives, for any possible distribution $\mathcal{U}$. A straightforward approach to is to introduce a hard constraint to the social distortion minimization (or DM maximization) problem, ensuring that at most a certain fraction of users are filtered out had they manipulated their content as much as they wished. More specifically, if a user $\mathbf{x}$ would like to follow a social trend $\mathbf{e}$ and move to the location $\mathbf{x} + \frac{\mathbf{e}}{2c}$, but by doing so their content gets filtered out, then their freedom of speech is violated. This leads to the following formalized problem:

$$\begin{aligned} & \text{find} && \arg \max_{f \in \mathcal{F}} \{ \mathbb{E}_{(\mathbf{x}, c) \sim \mathcal{U}} [h(f; \mathbf{x}, c)] \} \\ & \text{subject to} && \mathbb{E}_{(\mathbf{x}, c) \sim \mathcal{U}} [\mathbb{I}[f(\mathbf{x} + \frac{\mathbf{e}}{2c}) > 0]] \leq \theta. \end{aligned} \quad (7)$$

In reality, the platform usually only has access to an offline dataset $U = \{u = (\mathbf{x}_i, c_i)\}_{i=1}^n$ sampled from some distribution $\mathcal{U}$. Therefore, a practical way to estimate the solution of OP (7) is to solve the following empirical social distortion optimization problem. During training, we assume that we have access to un-manipulated examples. We can retrieve a set of clean examples by removing part of the content that violates the guidelines, e.g. misinformation.

$$\begin{aligned}
& \text{find} && \arg \max_{f \in \mathcal{F}} \left\{ \sum_{(\mathbf{x}_i, c_i) \in S} [h(f; \mathbf{x}, c)] \right\} \\
& \text{subject to} && \sum_{i=1}^n [\mathbb{I}[f(\mathbf{x}_i + \frac{\epsilon}{2c}) > 0]] \leq K.
\end{aligned} \tag{8}$$

In the following discussion, we will first examine how well the empirical solution to OP (8) approximates the solution to OP (7) using tools from standard statistical learning theory. We will then explore the computational aspects of solving OP (8).

4 LEARNABILITY OF GENERAL CONTENT MODERATORS

In this section we establish generalization guarantees for OP (7) – how many samples we need from the true distribution $\mathcal{D}$ to solve OP (8) to approximate the solution of OP (7). Theorem 1 shows sample complexity results in terms of the Vapnik–Chervonenkis dimension (VCDim) of the hypothesis moderator function class $\mathcal{F}$ and *Pseudo-Dimension* (PDim) of its corresponding distortion mitigation function class $\mathcal{H}_{\mathcal{F}}$.

Theorem 1. *For any moderator function class $\mathcal{F} = \{f : \mathbb{R}^d \rightarrow \{0, 1\}\}$ and its induced DM function class $\mathcal{H}_{\mathcal{F}} = \{h(f) | f \in \mathcal{F}\}$ defined by Eq. (5), and any distribution $\mathcal{U}$ on $\mathcal{X} \times \mathcal{C}$, a training sample U of size $\mathcal{O}\left(\frac{1}{\epsilon^2} \left(H^2(\text{PDim}(\mathcal{H}_{\mathcal{F}}) + \ln(1/\delta)) + \text{VCDim}(\mathcal{F})\right)\right)$ is sufficient to ensure that with probability at least $1 - \delta$, for every $f \in \mathcal{F}$, the distortion mitigation of f on U and $\mathcal{U}$ and the fraction of filtered points on U and $\mathcal{U}$ each differ by at most ϵ .*

Intuitively, Pseudo-dimension is a generalization of VC-dimension to real-valued function classes, capturing the capacity of a hypothesis class to fit continuous outputs rather than binary labels. Similar to VC-dimension, which measures the complexity of a class in terms of shattering points in binary classification, pseudo-dimension evaluates the ability of a function class to fit arbitrary real values over a set of points. The formal definition of Pseudo-dimension can be found in Appendix C.

Theorem 1 provides a general yet abstract characterization of the sample complexity required to approximate the solution to OP (7). More concretely, it suggests that problem (7) is statistically learnable if we focus on moderator function classes $\mathcal{F}$ with a finite VC-dimension and ensure that the corresponding class $\mathcal{H}$ has a finite Pseudo-Dimension. Fortunately, many natural function classes $\mathcal{F}$ satisfy these conditions, as shown in the following Proposition 3.

Proposition 3. *There exists function classes $\mathcal{F}$ such that $\text{VCDim}(\mathcal{F})$ and $\text{PDim}(\mathcal{H}_{\mathcal{F}})$ are both bounded, such as linear functions, piece-wise linear functions and kernel-based non-linear functions.*

Theorem 1, together with Proposition 3, demonstrates that finding the optimal linear moderator over an offline dataset for Eq. (8) is statistically efficient for many natural and practical function classes, including those discussed in Proposition 3. The linear class is arguably one of the simplest and most effective tools for moderation, capable of representing linear scoring rules that aggregate UGC scores based on relevant features. When combined with feature transformation mappings, linear models can represent techniques like dimensionality reduction followed by linear scoring. Many transformation techniques, such as invertible autoencoders, satisfy the invertibility requirement, ensuring statistical learnability. Additionally, piecewise linear function classes correspond to scenarios where multiple scoring rules are applied simultaneously. However, for such classes, the VCDim and PDim grow exponentially with the number of linear functions. Nevertheless, if the number of functions m remains small, sample-efficient learning is still achievable.

Proof of Theorem 1 (Appendix D) first applies standard learning theory for real-valued functions (Pollard, 1984) to establish a generalization bound for OP (8) without the freedom of speech constraint, relying on the PDim of $\mathcal{H}_{\mathcal{F}}$. Then, a union bound is used to account for the additional constraint, which depends on the VCDim of $\mathcal{F}$. The proof of Proposition 3 involves a detailed analysis of the closed-form best response

mapping for a user facing linear moderators. The core of the proof leverages the Sauer-Shelah Lemma (Sauer, 1972; Shelah, 1972) to establish an upper bound on the PDim for a composition of two function classes. Next, we study the computational complexity of empirically identifying a moderator that optimizes social distortion subject to freedom of speech constraints.

5 COMPUTATIONAL TRACTABILITY OF OPTIMAL LINEAR MODERATORS

We discuss the computational complexity of OP (8) in this section. To illustrate the idea, we focus on the class of linear moderator functions (i.e., $\mathcal{F} = \{f(\mathbf{x}) = \mathbf{w}^\top \mathbf{x} + b | \mathbf{w} \in \mathbb{R}^d, b \in \mathbb{R}\}$) as it yields a closed-form objective function, which makes the problem more tractable. And in order to also derive a closed-form for the constraint, we use the true feature $\mathbf{x}$ to filter content, but not the manipulated feature. Such an easier version of OP (8) is formulated by the following Lemma 1. And perhaps surprisingly, we show that this problem is NP-hard.

Lemma 1. *When $\mathcal{F} = \{f(\mathbf{x}) = \mathbf{w}^\top \mathbf{x} + b | \mathbf{w} \in \mathbb{R}^d, b \in \mathbb{R}\}$ is the linear function class, OP (8) is equivalent to the following constrained optimization problem:*

$$\begin{aligned} \text{find} \quad & \arg \min_{\mathbf{w}, b} \left\{ \sum_{i \in I} \left[(\mathbf{w}^\top \mathbf{x}_i + b)^2 - \left(\frac{\mathbf{w}^\top \mathbf{e}}{2c_i} \right)^2 \right] \right\} \\ \text{subject to} \quad & \sum_{i=1}^n \mathbb{I} \left[\mathbf{w}^\top \left(\mathbf{x}_i + \frac{\mathbf{e}}{2c_i} \right) + b \leq 0 \right] \geq n - K, \\ & -1 \leq w_j \leq 1, 1 \leq j \leq d. \end{aligned} \quad (9)$$

where the index set $I = \{i \in [n] : -\frac{\mathbf{w}^\top \mathbf{e}}{2c_i} < \mathbf{w}^\top \mathbf{x}_i + b \leq 0\}$.

Since OP (9) involves the strict constraint $-\frac{\mathbf{w}^\top \mathbf{e}}{2c_i} < \mathbf{w}^\top \mathbf{x}_i + b$, we follow a standard practice by introducing a slack variable $\epsilon > 0$ and consider a relaxed problem, replacing the strict constraint with a non-strict one: $\epsilon - \frac{\mathbf{w}^\top \mathbf{e}}{2c_i} \leq \mathbf{w}^\top \mathbf{x}_i + b$. A natural question that follows is whether we can efficiently solve this relaxed version of OP (9). However, despite the nice quadratic form of the objective function in OP (9), the combinatorial nature of the constraint and the indefiniteness of the quadratic objective make the problem challenging to solve. In fact, the next Theorem 2 shows that any ϵ -relaxation of OP (9) is NP-hard.

Theorem 2. *For any given input $\epsilon > 0, n, K \in \mathbb{N}_+$ and offline dataset $\mathcal{X} = \{(\mathbf{x}_i, c_i)\}_{i=1}^n$, finding the optimal solution to the ϵ -relaxation of OP (9) (by replacing the index set I with $I_\epsilon = \{i \in [n] : \epsilon - \frac{\mathbf{w}^\top \mathbf{e}}{2c_i} < \mathbf{w}^\top \mathbf{x}_i + b \leq 0\}$) is NP-hard with respect to $(n, K, 1/\epsilon)$.*

Theorem 2 demonstrates that minimizing social distortion under a hard constraint—limiting the number of users whose content can be filtered—is computationally intractable. This complexity arises because finding a linear moderator f that minimizes social distortion is analogous to finding a hyperplane that maximizes the number of points near its boundary, as the amount of social distortion for each content $\mathbf{x}$ only increases as $\mathbf{x}$ approaches the boundary of f . With the additional constraint, the problem becomes a combinatorial geometric challenge: given a set of n points, find a hyperplane that maximizes the number of points lying on it while ensuring that at least K points remain on each side of the hyperplane. This turns out to be hard. The formal proof, provided in Appendix E.2, contains two core reductions. First, we reduce OP (9) from a combinatorial optimization problem called Maximum Feasible Linear Subsystems (MAX-FLS) with mandatory constraints, and then we show that the problem of MAX-FLS with mandatory constraints is NP-hard by showing a reduction from the Exact 3-Set Cover problem, which is known to be NP-hard.

5.1 A HEURISTIC METHOD FOR FINDING APPROXIMATED SOLUTIONS

Since minimizing the social distortion with a hard freedom of speech constraint is NP-hard even for linear function class, we resort to an approximation approach for solving this problem. Still focusing on linear

moderators, a straightforward method for tackling the hard constraint is using penalty functions (Han & Mangasarian, 1979; Nocedal & Wright, 1999). That is, for any $(\mathbf{x}_i, c_i)$ that violates the moderator, we introduce a penalty function $P_i(\mathbf{w}, b)$ in the objective, as formulated in the following OP:

$$\begin{aligned} \arg \min_{\mathbf{w}, b} & \left\{ \sum_{i \in I} \left[(\mathbf{w}^\top \mathbf{x}_i + b)^2 - \left(\frac{\mathbf{w}^\top \mathbf{e}}{2c_i} \right)^2 \right] + \lambda P(\mathbf{w}, b) \right\} \\ \text{s.t.} & -1 \leq w_j \leq 1, 1 \leq j \leq d. \end{aligned} \quad (10)$$

Let $g(\mathbf{w}, b; U, \mathbf{e}) = \left[\mathbf{w}^\top \left(\mathbf{x}_i + \frac{\mathbf{e}}{2c_i} \right) + b > 0 \right]_+$, where the notation $[y_i]_+$ denotes the vector $[\max(0, y_i)]_{i=1}^n$. The hard constraint in OP (9) can thus be expressed as $\|g(\mathbf{w}, b; U, \mathbf{e})\|_0 \leq K$, where $\|\cdot\|_0$ represents the ℓ_0 -norm. A natural choice for the penalty in OP (10) is then $P(\mathbf{w}, b) = \|g(\mathbf{w}, b; U, \mathbf{e})\|_0$. However, in machine learning and optimization, it is common practice to replace the ℓ_0 penalty with ℓ_1 - or ℓ_2 -based penalties, often referred to as convex relaxations, to improve numerical stability. We will use the ℓ_2 -based penalty in our experiments.

Finding the optimal penalty strength λ : To balance the social distortion objective and the freedom of speech penalty term, selecting an appropriate parameter $\lambda > 0$ is crucial. A value of λ that is too small may fail to yield a solution satisfying the constraint, whereas a sufficiently large λ can enforce the constraint but may lead to an excessively large objective value (i.e., a too small DM). Intuitively, there exists an optimal choice of λ that strikes this balance. The following proposition demonstrates the existence of such a λ and outlines how the platform can determine it.

Proposition 4. *Let $(\mathbf{w}_1, b_1)$, $(\mathbf{w}_2, b_2)$ be any solution of OP (10) by choosing $\lambda = \lambda_1$ and $\lambda = \lambda_2$ ($\lambda_1 > \lambda_2 \geq 0$). Then we have $P(\mathbf{w}_1, b_1) \leq P(\mathbf{w}_2, b_2)$, $DM(\mathbf{w}_1, b_1) \leq DM(\mathbf{w}_2, b_2)$.*

Proposition 4 enables the platform to employ a binary search-based approach to determine a proper λ for any given constraint $\|g(\mathbf{w}, b)\|_0 \leq K$. First, the platform can start with $\lambda_{\max} = \sup_{(\mathbf{w}, b)} \{DM(\mathbf{w}, b)\} \sim O(n^3)$, as any $\lambda > \lambda_{\max}$ would enforce the equivalent constraint $\|g(\mathbf{w}, b)\|_0 = 0$. Then, the platform performs a binary search in interval $[\lambda_{\text{left}}, \lambda_{\text{right}}] = [0, \lambda_{\max}]$. At each iteration, it solves OP (10) with $\lambda = \frac{\lambda_{\text{left}} + \lambda_{\text{right}}}{2}$ and obtains $(\mathbf{w}_*, b_*)$. According to Proposition 4, if $P(\mathbf{w}_*, b_*) > K$, the platform excludes the lower half of the range $[\lambda_{\text{left}}, \lambda_{\text{right}}]$ in the next round; otherwise, it searches in the upper half. This procedure allows the platform to identify the optimal λ within precision δ by solving at most $O(3 \log_2(n) + \log_2(1/\delta))$ instances of OP (10). The proof of Proposition 4 is straightforward and is provided in Appendix E.3.

Efficient heuristic with smooth surrogate loss: Next we discuss for a particular chosen λ , how to efficiently obtain an solution of OP (10). Let $a_i = \frac{\mathbf{w}^\top \mathbf{e}}{2c_i}$, $y_i = \mathbf{w}^\top \mathbf{x}_i + b + a_i$, we can re-formulate OP (10) as the following cleaner form

$$\begin{aligned} \text{find} & \quad \arg \min_{\mathbf{w}, b} \left\{ \sum_{1 \leq i \leq n} l(\mathbf{w}, b; \mathbf{x}_i, c_i, \mathbf{e}) \right\} \\ \text{subject to} & \quad l(\mathbf{w}, b; \mathbf{x}_i, c_i, \mathbf{e}) = \max \{0, y_i\} \cdot (y_i - 2a_i) + \lambda P_i(\mathbf{w}, b) \\ & \quad -1 \leq w_j \leq 1, 1 \leq j \leq d. \end{aligned} \quad (11)$$

The objective function in OP (11) can be understood as the aggregation of the social good loss l induced by each user i , consisting of two components. The first part, $\max \{0, y_i\} \cdot (y_i - 2a_i)$, measures the social distortion incurred by the linear moderator $(\mathbf{w}, b)$, and the second part reflects the calibrated infringement on freedom of speech: the larger penalty term P_i is, the farther user i 's content is from the decision boundary on the positive side of f , making it more likely that user i 's content $\mathbf{x}$ will be filtered.

The structure of OP (11) resembles the empirical loss minimization problem commonly seen in standard machine learning problems, and we can employ a stochastic gradient descent approach to tackle it, given

any specific penalty functions and trade-off parameter λ . To ensure the social good loss l is differentiable so that we can apply gradient-based approach, we need to further introduce a surrogate loss $\tilde{l}$ to smooth the non-differentiable point at $y_i = 0$ of $\max\{0, y_i\} \cdot (y_i - 2a_i)$ while selecting a differentiable penalty function. The detail of this treatment is similar to (Levanon & Rosenfeld, 2021) and is outlined in the optimization solver setup in the next section. In the following experiments, we apply this approach to solve (11) using a synthetic dataset and report the approximate optimal linear moderator for different trade-off parameters λ .

Experiment Result: We use a synthetic dataset and perform experiments of our heuristic optimization method. The details of the dataset and the setup of optimizers can be found in Appendix F.1. A 2-dimensional visualization in Figure 3 illustrates the optimal linear moderators for both a small λ ($\lambda = 0.1$) and a larger value ($\lambda = 10.0$). Each user’s original content feature $\mathbf{x}$ is represented by a blue dot, while its strategic response to the moderator is shown in red. The social trend is $\mathbf{e} = (1, 0)$. As the figure shows, a larger λ shifts the moderator boundary toward the margin of the content distribution, as desired. This results in fewer pieces of content being filtered, while still achieving a reasonable degree of social distortion mitigation. When λ is small, the computed optimal moderator minimizes social distortion but at the expense of infringing on more users’ freedom of speech. The right panel displays both the social distortion mitigation (i.e., the negative of the optimal objective value of OP (11)) and a freedom of speech preservation index, measured by the fraction of content that remains on the platform under the regulation of the computed moderators with varying λ . As shown, the freedom of speech index increases as λ grows, while social distortion mitigation follows an inverted U-shape. This suggests a trade-off between these two objectives, similar to the one observed in the toy model 2. Our result indicates our proposed optimization technique can effectively approximate a solution, thus allowing the platform to flexibly balance social distortion and freedom of speech despite of the computational challenge.

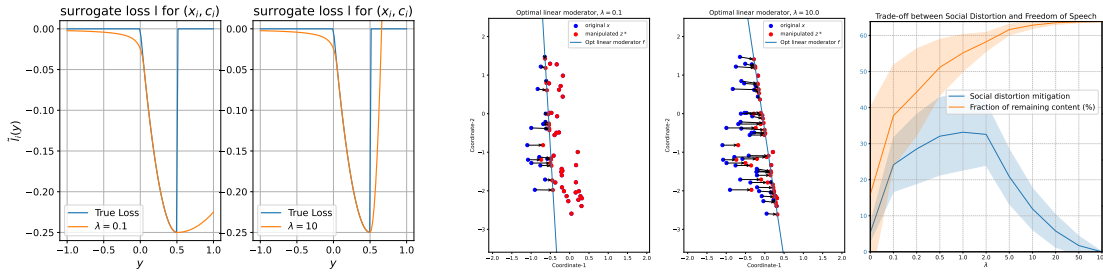


Figure 3: Left: the constructed quasi-convex single point surrogate loss function. Middle: the computed moderator obtained under $\lambda = 0.1$ and 10.0 . Arrows represent users’ strategic manipulations against the optimal linear moderator. Right: social distortion mitigation (blue) and the fraction of remaining content on the platform (yellow) incurred by the computed moderator obtained under different $\lambda \in [0.1, 100]$. Error bars obtained from results with 20 independently generated dataset. Error bars are 1σ region based on results from 20 independently generated datasets.

6 CONCLUSION

We addressed the challenge of designing content moderators that reduce engagement with harmful social trends while preserving freedom of speech. By modeling the problem as a constrained optimization, we introduced the concept of social distortion and provided generalization guarantees based on the VC-dimension and Pseudo-dimension of the filter function class. While we established the computational hardness of finding optimal linear filters, we provide an empirically efficient approximation approach that enables the platform to achieve any desirable trade-offs. Our findings highlight the need for efficient algorithms and further exploration of more flexible filtering mechanisms.

REFERENCES

- Saba Ahmadi, Hedyeh Beyhaghi, Avrim Blum, and Keziah Naggita. The strategic perceptron. pp. 6–25, 2021.
- Saba Ahmadi, Hedyeh Beyhaghi, Avrim Blum, and Keziah Naggita. On classification of strategic agents who can both game and improve. *arXiv preprint arXiv:2203.00124*, 2022.
- Tal Alon, Magdalen Dobson, Ariel Procaccia, Inbal Talgam-Cohen, and Jamie Tucker-Foltz. Multiagent evaluation mechanisms. 34(02):1774–1781, 2020.
- Edoardo Amaldi and Viggo Kann. The complexity and approximability of finding maximum feasible subsystems of linear relations. *Theoretical computer science*, 147(1-2):181–210, 1995.
- Martin Anthony, Peter L Bartlett, Peter L Bartlett, et al. *Neural network learning: Theoretical foundations*, volume 9. cambridge university press Cambridge, 1999.
- Solon Barocas and Andrew D Selbst. Big data’s disparate impact. *Calif. L. Rev.*, 104:671, 2016.
- Yahav Bechavod, Katrina Ligett, Steven Wu, and Juba Ziani. Gaming helps! learning from strategic interactions in natural dynamics. In *International Conference on Artificial Intelligence and Statistics*, pp. 1234–1242. PMLR, 2021.
- Alessandro Bessi and Emilio Ferrara. Social bots distort the 2016 us presidential election online discussion. *First monday*, 21(11-7), 2016.
- Olga Boichak, Sam Jackson, Jeff Hemsley, and Sikana Tanupabrungsun. Automated diffusion? bots and their influence during the 2016 us presidential election. In *Transforming Digital Worlds: 13th International Conference, iConference 2018, Sheffield, UK, March 25-28, 2018, Proceedings 13*, pp. 17–26. Springer, 2018.
- Lee C Bollinger and Geoffrey R Stone. *Social media, freedom of speech, and the future of our democracy*. Oxford University Press, 2022.
- Valerie C Brannon. Free speech and the regulation of social media content. *Congressional Research Service (CRS) Reports and Issue Briefs*, pp. NA–NA, 2019.
- Michael Brückner and Tobias Scheffer. Stackelberg games for adversarial prediction problems. In *Proceedings of the 17th ACM SIGKDD international conference on Knowledge discovery and data mining*, pp. 547–555, 2011.
- Caroline Cauffman and Catalina Goanta. A new order: The digital services act and consumer protection. *European Journal of Risk Regulation*, 12(4):758–774, 2021.
- Yiling Chen, Yang Liu, and Chara Podimata. Learning strategy-aware linear classifiers. volume 33, pp. 15265–15276, 2020.
- Nilesh Dalvi, Pedro Domingos, Mausam, Sumit Sanghai, and Deepak Verma. Adversarial classification. In *Proceedings of the tenth ACM SIGKDD international conference on Knowledge discovery and data mining*, pp. 99–108, 2004.
- Ofer Dekel, Felix Fischer, and Ariel D Procaccia. Incentive compatible regression learning. In *Proceedings of the nineteenth annual ACM-SIAM symposium on Discrete algorithms*, pp. 884–893, 2008.
- Jinshuo Dong, Aaron Roth, Zachary Schutzman, Bo Waggoner, and Zhiwei Steven Wu. Strategic classification from revealed preferences. pp. 55–70, 2018.

- Ysabel Gerrard. Beyond the hashtag: Circumventing content moderation on social media. *New Media & Society*, 20(12):4492–4511, 2018.
- Tarleton Gillespie. *Custodians of the Internet: Platforms, content moderation, and the hidden decisions that shape social media*. Yale University Press, 2018.
- Nika Haghtalab, Nicole Immorlica, Brendan Lucier, and Jack Z. Wang. Maximizing welfare with incentive-aware evaluation mechanisms. pp. 160–166, 2020.
- S P Han and Olvi L Mangasarian. Exact penalty functions in nonlinear programming. *Mathematical programming*, 17:251–269, 1979.
- Moritz Hardt, Nimrod Megiddo, Christos Papadimitriou, and Mary Wootters. Strategic classification. pp. 111–122, 2016.
- Lily Hu, Nicole Immorlica, and Jennifer Wortman Vaughan. The disparate effects of strategic manipulation. In *Proceedings of the Conference on Fairness, Accountability, and Transparency*, FAT* ’19, pp. 259–268, 2019.
- Jon Kleinberg and Manish Raghavan. How do classifiers induce agents to invest effort strategically? *ACM Transactions on Economics and Computation (TEAC)*, 8(4):1–23, 2020.
- Kate Klonick. The new governors: The people, rules, and processes governing online speech. *Harv. L. Rev.*, 131:1598, 2017.
- Sagi Levanon and Nir Rosenfeld. Strategic classification made practical. In *International Conference on Machine Learning*, pp. 6243–6253. PMLR, 2021.
- MURPHY COLIN MICHAEL. Cyberbullying among young people: Laws and policies in selected member states. 2024.
- Smitha Milli, John Miller, Anca D Dragan, and Moritz Hardt. The social cost of strategic classification. In *Proceedings of the Conference on Fairness, Accountability, and Transparency*, pp. 230–239, 2019.
- Jorge Nocedal and Stephen J Wright. *Numerical optimization*. Springer, 1999.
- Catherine O’Regan and Stefan Theil. Hate speech regulation on social media: An intractable contemporary challenge. *Research Outreach*, 2020.
- Juan Perdomo, Tijana Zrnic, Celestine Mendler-Dünner, and Moritz Hardt. Performative prediction. volume 119, pp. 7599–7609. PMLR, 2020.
- D. Pollard. *Convergence of Stochastic Processes*. Springer New York, 1984. ISBN 9780387909905. URL <https://books.google.com/books?id=B2vgGMa9vd4C>.
- Sarah T Roberts. Digital detritus: ‘error’ and the logic of opacity in social media content moderation. *First Monday*, 2018.
- Nir Rosenfeld and Haifeng Xu. Machine learning should maximize welfare, not (only) accuracy. *arXiv preprint arXiv:2502.11981*, 2025.
- Norbert Sauer. On the density of families of sets. *Journal of Combinatorial Theory, Series A*, 13(1):145–147, 1972.

- Brittany Shaughnessy, Eliana DuBosar, Myiah J Hutchens, and Ilyssa Mann. An attack on free speech? examining content moderation,(de-), and (re-) platforming on american right-wing alternative social media. *New Media & Society*, pp. 14614448241228850, 2024.
- Saharon Shelah. A combinatorial problem; stability and order for models and theories in infinitary languages. *Pacific Journal of Mathematics*, 41(1):247–261, 1972.
- Ravi Sundaram, Anil Vullikanti, Haifeng Xu, and Fan Yao. Pac-learning for strategic classification. *Journal of Machine Learning Research*, 24(192):1–38, 2023.
- Jessica West. *Cyber-violence against women*. Battered Women’s Support Services, 2014.
- Fan Yao, Yiming Liao, Jingzhou Liu, Shaoliang Nie, Qifan Wang, Haifeng Xu, and Hongning Wang. Unveiling user satisfaction and creator productivity trade-offs in recommendation platforms. *Advances in Neural Information Processing Systems*, 37:86958–86984, 2024a.
- Fan Yao, Yiming Liao, Mingzhe Wu, Chuanhao Li, Yan Zhu, James Yang, Jingzhou Liu, Qifan Wang, Haifeng Xu, and Hongning Wang. User welfare optimization in recommender systems with competing content creators. In *Proceedings of the 30th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*, pp. 3874–3885, 2024b.
- Greyson K Young. How much is too much: the difficulties of social media content moderation. *Information & Communications Technology Law*, 31(1):1–16, 2022.

Appendix to *Strategic Filtering for Content Moderation: Free Speech or Free of Distortion?*

A ADDITIONAL RELATED WORK

Welfare-oriented learning. The key novelty of our work is a shift from mere predictive performance to welfare-oriented optimization in strategic learning, aligning with a recent call by Rosenfeld & Xu (2025). The goal of standard strategic classification tasks is to preserve accuracy despite agents’ manipulations while in our setting accuracy is not well-defined as there is no ground-truth labels for UGC. What we advocate is that the platform should directly optimize a welfare objective that captures the externalities of moderation decisions targeting harmful social trends. Our work also connects with a broader literature that examines the societal impacts of strategic classification. For example, Milli et al. (2019) highlights how counter-strategic defenses increase social burden, Kleinberg & Raghavan (2020) explores how classifiers can steer gaming behavior toward socially beneficial efforts, and Yao et al. (2024b;a) studies how recommender systems can improve user welfare in the presence of strategic content creators. We extend this perspective to a new and timely domain — balancing social distortion and freedom of speech in content moderation.

Content moderation in social media platforms. To detect abusive content, platforms use a mix of human moderators and automated algorithms. While early moderation relied heavily on human review (Klonick, 2017), most platforms now employ automated filters to handle overtly inappropriate content (Gillespie, 2018), with humans in the loop to address limitations such as poor generalization to out-of-distribution cases. In this work, we assume the harmful trend is known, e.g., election misinformation, and focus on designing mechanisms that discourage user engagement with such trends while preserving freedom of speech.

B PROOF OF PROPOSITION 1

Proof. Let $\mathcal{D} = \{z : f(z; w) \leq 0\}$. According to the definition of convex moderator, $\mathcal{D}$ is a convex set in $\mathbb{R}^d$.

If $x \in \mathcal{D}$, under the formulation of Eq. (1), each user’s best response is the solution to the following convex optimization problem (OP)

$$\begin{aligned} \text{find} \quad & z^* = \arg \min_z \{-z^\top e + c\|z - x\|_2^2\} \\ \text{subject to} \quad & z \in \mathcal{D}. \end{aligned} \tag{12}$$

Observe that the objective function

$$-z^\top e + c\|z - x\|_2^2 = c\left\|z - \left(x + \frac{e}{2c}\right)\right\|_2^2 - \frac{e^\top e}{4c},$$

OP (12) is thus equivalent to

$$\begin{aligned} \text{find} \quad & z^* = \arg \min_z \left\{\left\|z - \left(x + \frac{e}{2c}\right)\right\|_2^2\right\} \\ \text{subject to} \quad & z \in \mathcal{D}. \end{aligned} \tag{13}$$

Let $z' = x + \frac{e}{2c}$. If z' is feasible, i.e., $f(z') \leq 0$, we have $z^* = z'$. Otherwise, by definition z^* is the ℓ_2 projection of z' on to the decision boundary of f .

If $x \notin \mathcal{D}$, staying at x yield 0 utility for u . As a result, $z^* = \mathcal{P}_f(z')$ only when z' yields a negative objective value in OP (12). Otherwise, $z^* = x$.

□

C DEFINITION OF PSEUDO-DIMENSION

Definition 3. (Pollard’s Pseudo-Dimension) A class $\mathcal{F}$ of real-valued functions P -shatters a set of points $\mathcal{X} = \{x_1, x_2, \dots, x_n\}$ if there exists a set of thresholds $\gamma_1, \gamma_2, \dots, \gamma_n$ such that for every subset $T \subseteq \mathcal{X}$, there exists a function $f_T \in \mathcal{F}$ such that $f_T(x_i) \geq \gamma_i$ if and only if $x_i \in T$. In other words, all 2^n possible above/below patterns are achievable for targets $\gamma_1, \dots, \gamma_n$. The pseudo-dimension of $\mathcal{F}$, denoted by $\text{PDim}(\mathcal{F})$, is the size of the largest set of points that it P -shatters.

D OMITTED PROOFS IN SECTION 4

D.1 PROOF OF PROPOSITION 3

In this section we present the proof of Proposition 3, which upper bounds the PDim of the distortion mitigation (DM) function class $\mathcal{H}_{\mathcal{F}}$ given some example moderator function class $\mathcal{F}$. Specifically, we will show that

1. When $\mathcal{F} = \{f(\mathbf{x}) = \mathbb{I}[\mathbf{w}^\top \mathbf{x} + b \leq 0] | (\mathbf{w}, b) \in \mathbb{R}^{d+1}\}$ is the linear class, we have

$$\text{VCDim}(\mathcal{F}) \leq d + 1, \quad \text{PDim}(\mathcal{H}_{\mathcal{F}}) \leq \tilde{\mathcal{O}}(d^2), \quad (14)$$

where $\tilde{\mathcal{O}}$ is the big $\mathcal{O}$ notation omitting the log terms.

2. When $\mathcal{F}$ is a piece-wise linear function class with each instance constitutes m linear functions, i.e., $\mathcal{F} = \{f(\mathbf{x}) = \mathbb{I}[\mathbf{w}_1^\top \mathbf{x} + b_1 \leq 0] \vee \dots \vee \mathbb{I}[\mathbf{w}_m^\top \mathbf{x} + b_m \leq 0] | (\mathbf{w}_i, b_i) \in \mathbb{R}^{d+1}, 1 \leq i \leq m\}$, we have

$$\text{VCDim}(\mathcal{F}) \leq \tilde{\mathcal{O}}(d \cdot 3^m), \quad \text{PDim}(\mathcal{H}_{\mathcal{F}}) \leq \tilde{\mathcal{O}}(d^{m+1} \cdot 3^{2^m}). \quad (15)$$

3. When $\mathcal{F}$ is the linear class defined on some feature transformation mapping ϕ , i.e., $\mathcal{F} = \{f(\mathbf{x}) = \mathbb{I}[\mathbf{w}^\top \phi(\mathbf{x}) + b \leq 0] | (\mathbf{w}, b) \in \mathbb{R}^{d+1}\}$, as long as ϕ is invertible and order-preserving, it also holds that

$$\text{VCDim}(\mathcal{F}) \leq d + 1, \quad \text{PDim}(\mathcal{H}_{\mathcal{F}}) \leq \tilde{\mathcal{O}}(d^2). \quad (16)$$

Proof. In this proof we derive Pseudo-dimension upper bounds for the three cases listed.

Case-1: When $\mathcal{F} = \{f(\mathbf{x}) = \mathbb{I}[\mathbf{w}^\top \mathbf{x} + b \leq 0] | (\mathbf{w}, b) \in \mathbb{R}^{d+1}\}$ is the linear functions class, we can without loss of generality let $\|\mathbf{w}\|_2 = 1$ since a simultaneous rescaling of $\mathbf{w}, b$ does not change the nature of the moderator function and its induced strategic responses. Next, we derive the DM class $\mathcal{H}_{\mathcal{F}}$ as follows. First of all, plugging in the expression of f into the result of Proposition 1, we obtain a user $(\mathbf{x}, c)$ ’s best response as the following:

1. if $\mathbf{w}^\top \cdot (\mathbf{x} + \frac{\mathbf{e}}{2c}) + b \leq 0, \mathbf{z}^* = \mathbf{x} + \frac{\mathbf{e}}{2c}$.
2. if $\mathbf{w}^\top \cdot (\mathbf{x} + \frac{\mathbf{e}}{2c}) + b > 0$ and $\mathbf{w}^\top \mathbf{x} + b \leq 0, \mathbf{z}^* = P_{\mathcal{F}}(\mathbf{x} + \frac{\mathbf{e}}{2c})$ which has the following closed-form expression

$$\begin{aligned} \mathbf{z}^* &= \mathbf{x} + \frac{\mathbf{e}}{2c} - \frac{\mathbf{w}^\top (\mathbf{x} + \frac{\mathbf{e}}{2c}) \mathbf{w}}{\mathbf{w}^\top \mathbf{w}} - \frac{b \mathbf{w}}{\mathbf{w}^\top \mathbf{w}} \\ &= \mathbf{x} + \frac{\mathbf{e}}{2c} - \mathbf{w}^\top (\mathbf{x} + \frac{\mathbf{e}}{2c}) \mathbf{w} - b \mathbf{w}. \end{aligned} \quad (17)$$

By Definition 1 and 2, we can compute each function $h \in \mathcal{H}_{\mathcal{F}}$ as

$$\begin{aligned} h(f, e; \mathbf{x}, c) &= D(\perp; (\mathbf{x}, c), e) - D(f; (\mathbf{x}, c), e) \\ &= \mathbb{I}[\mathbf{w}^\top \mathbf{x} + b \leq 0] \cdot \mathbb{I}\left[\mathbf{w}^\top \mathbf{x} + b > -\frac{\mathbf{w}^\top \mathbf{e}}{2c}\right] \cdot \left(\left\|\frac{\mathbf{e}}{2c}\right\|_2^2 - \left\|\frac{\mathbf{e}}{2c} - \mathbf{w}^\top \left(\mathbf{x} + \frac{\mathbf{e}}{2c}\right) \mathbf{w} - b \mathbf{w}\right\|_2^2\right) \\ &= \mathbb{I}[\mathbf{w}^\top \mathbf{x} + b \leq 0] \cdot \mathbb{I}\left[\mathbf{w}^\top \mathbf{x} + b > -\frac{\mathbf{w}^\top \mathbf{e}}{2c}\right] \cdot \left[-(\mathbf{w}^\top \mathbf{x} + b)^2 + \left(\frac{\mathbf{w}^\top \mathbf{e}}{2c}\right)^2\right]. \end{aligned} \quad (18)$$

For the ease of notation, let's define $\tilde{\mathbf{x}} = (\mathbf{x}, \frac{1}{2c}) \in \mathbb{R}^{d+1}$ be the extended feature vector for any user data $(\mathbf{x}, c)$. By the definition of Pseudo-dimension, for any function class $\mathcal{F} = \{f(\tilde{\mathbf{x}}; \mathbf{w}, b) | \mathbf{w}, b\}$, the $PD\dim(\mathcal{F})$ can be reduced to the VC dimension of the epigraph of $\mathcal{F}$, i.e.,

$$PD\dim(\mathcal{F}) = VC\dim(\{h(\tilde{\mathbf{x}}, y) = \text{sgn}(f(\tilde{\mathbf{x}}) - y) | f \in \mathcal{F}, y \in [-1, 1]\}). \quad (19)$$

Let's define the following three function classes

$$\begin{aligned} \mathcal{H}_1(d) &= \left\{ h_1(\mathbf{x}, c; \mathbf{w}, b) = -(\mathbf{w}^\top \mathbf{x} + b)^2 + \frac{(\mathbf{w}^\top \mathbf{e})^2}{4c^2} \mid (\mathbf{x}, c) \in \mathbb{R}^{d+1}, (\mathbf{w}, b) \in \mathbb{R}^{d+1} \right\} \\ &= \left\{ h_1(\tilde{\mathbf{x}}; \mathbf{w}, b) = -(\mathbf{w}^\top \tilde{\mathbf{x}}_{1:d} + b)^2 + (\mathbf{w}^\top \mathbf{e})^2 \tilde{x}_{d+1}^2 \mid \tilde{\mathbf{x}} \in \mathbb{R}^{d+1}, (\mathbf{w}, b) \in \mathbb{R}^{d+1} \right\}, \\ \mathcal{H}_2(d) &= \left\{ h_2(\mathbf{x}, c; \mathbf{w}, b) = \mathbb{I}\left[\mathbf{w}^\top \left(\mathbf{x} + \frac{\mathbf{e}}{2c}\right) + b \geq 0\right] \mid (\mathbf{x}, c) \in \mathbb{R}^{d+1}, (\mathbf{w}, b) \in \mathbb{R}^{d+1} \right\} \\ &= \left\{ h_2(\tilde{\mathbf{x}}; \mathbf{w}, b) = \mathbb{I}\left[\mathbf{w}^\top \tilde{\mathbf{x}}_{1:d} + (\mathbf{w}^\top \mathbf{e}) \tilde{x}_{d+1} + b \geq 0\right] \mid \tilde{\mathbf{x}} \in \mathbb{R}^{d+1}, (\mathbf{w}, b) \in \mathbb{R}^{d+1} \right\}, \\ \mathcal{H}_3(d) &= \left\{ h_3(\mathbf{x}; \mathbf{w}, b) = \mathbb{I}[\mathbf{w}^\top \mathbf{x} + b \leq 0] \mid \mathbf{x} \in \mathbb{R}^d, (\mathbf{w}, b) \in \mathbb{R}^{d+1} \right\}, \end{aligned}$$

where $\mathbf{x}_{1:d}$ denotes the vector that contains the first d -dimension of $\mathbf{x}$.

Since $h(\mathbf{w}, b; \tilde{\mathbf{x}}) = h_1 * h_2 * h_3(\tilde{\mathbf{x}}; \mathbf{w}, b)$, the Pseudo-dimension of $\mathcal{H}_{\mathcal{F}}$ can be upper bounded by the following

$$PD\dim(\mathcal{H}_{\mathcal{F}}) \leq VC\dim\left(\left\{h(\tilde{\mathbf{x}}, y) = \text{sgn}\left(\prod_{i=1}^3 h_i(\tilde{\mathbf{x}}) - y\right) \mid h_i \in \mathcal{H}_i, y \in [-1, 1], 1 \leq i \leq 3\right\}\right), \quad (20)$$

where the inequality holds because

$$\left\{\text{sgn}(h_1 * h_2 * h_3(\tilde{\mathbf{x}}; \mathbf{w}, b) - y) \mid (\mathbf{w}, b) \in \mathbb{R}^{d+1}\right\} \subseteq \left\{\text{sgn}\left(\prod_{i=1}^3 h_i(\tilde{\mathbf{x}}) - y\right) \mid h_i \in \mathcal{H}_i, 1 \leq i \leq 3\right\}.$$

For any function classes $\mathcal{F}, \mathcal{G}$, define $\mathcal{F} \otimes \mathcal{G} = \{f * g \mid f \in \mathcal{F}, g \in \mathcal{G}\}$. Then Eq. (20) suggests that in order to upper bound $PD\dim(\mathcal{H}_{\mathcal{F}})$, it suffices to upper bound $PD\dim(\mathcal{H}_1 \otimes \mathcal{H}_2 \otimes \mathcal{H}_3)$. Thanks to Lemma 2 which establishes the $PD\dim$ of the product of two function classes, this can be done by upper bounding $PD\dim(\mathcal{H}_1), PD\dim(\mathcal{H}_2), PD\dim(\mathcal{H}_3)$ separately. In the following we derive the $PD\dim$ for each function class $\mathcal{H}_i, i = 1, 2, 3$ and then use Lemma 2 to conclude the proof.

Deriving $PD\dim(\mathcal{H}_3)$: First of all, since the Pseudo-Dimension for a binary value function class is exactly the VC dimension of the corresponding real-valued function class inside the indicator function, we immediately obtain

$$PD\dim(\mathcal{H}_3) \leq d + 1, \quad (21)$$

which is the VC dimension for a d dimensional linear function class.

Deriving $PDim(\mathcal{H}_2)$: For $\mathcal{H}_2$, it holds that

$$\begin{aligned}\mathcal{H}_2(d) &= \left\{ h_2(\tilde{\mathbf{x}}; \mathbf{w}, b) = \mathbb{I}[\mathbf{w}^\top \tilde{\mathbf{x}}_{1:d} + (\mathbf{w}^\top \mathbf{e}) \tilde{x}_{d+1} + b \geq 0] \mid \tilde{\mathbf{x}} \in \mathbb{R}^{d+1}, (\mathbf{w}, b) \in \mathbb{R}^{d+1} \right\} \\ &\subset \left\{ h_2(\tilde{\mathbf{x}}; \mathbf{w}, w_{d+1}, b) = \mathbb{I}[\mathbf{w}^\top \tilde{\mathbf{x}}_{1:d} + w_{d+1} \tilde{x}_{d+1} + b \geq 0] \mid \tilde{\mathbf{x}} \in \mathbb{R}^{d+1}, (\mathbf{w}, w_{d+1}, b) \in \mathbb{R}^{d+2} \right\} \\ &= \left\{ h_2(\tilde{\mathbf{x}}; \tilde{\mathbf{w}}, b) = \mathbb{I}[\tilde{\mathbf{w}}^\top \tilde{\mathbf{x}} + b \geq 0] \mid \tilde{\mathbf{x}} \in \mathbb{R}^{d+1}, (\tilde{\mathbf{w}}, b) \in \mathbb{R}^{d+2} \right\},\end{aligned}\quad (22)$$

where the subset relationship (22) holds because we relax the parameter $\mathbf{w}^\top \mathbf{e}$ correlated with $\mathbf{w}$ to an additional independent parameter $w_{d+1} \in \mathbb{R}$. This implies that $\mathcal{H}_2(d)$ is a subclass of indicator functions induced by the $d+1$ dimensional linear class. As a result, the Pseudo-Dimension of $\mathcal{H}_2$ must be upper bounded by $d+2$.

Deriving $PDim(\mathcal{H}_1)$: To derive $PDim(\mathcal{H}_1)$, we first apply the same trick to relax $\mathbf{w}^\top \mathbf{e}$ to an independent parameter w_{d+1} :

$$\begin{aligned}\mathcal{H}_1(d) &= \left\{ h_1(\tilde{\mathbf{x}}; \mathbf{w}, b) = (\mathbf{w}^\top \tilde{\mathbf{x}}_{1:d} + b)^2 - (\mathbf{w}^\top \mathbf{e})^2 \tilde{x}_{d+1}^2 \mid \tilde{\mathbf{x}} \in \mathbb{R}^{d+1}, (\mathbf{w}, b) \in \mathbb{R}^{d+1} \right\} \\ &\subset \left\{ h_1(\tilde{\mathbf{x}}; \tilde{\mathbf{w}}, b) = (\mathbf{w}^\top \tilde{\mathbf{x}}_{1:d} + b)^2 - w_{d+1}^2 \tilde{x}_{d+1}^2 \mid \tilde{\mathbf{x}} \in \mathbb{R}^{d+1}, (\tilde{\mathbf{w}}, b) \in \mathbb{R}^{d+2} \right\} \triangleq \tilde{\mathcal{H}}_1(d),\end{aligned}\quad (23)$$

where $\tilde{\mathbf{w}} = (\mathbf{w}, w_{d+1})$. Note that each instance in $\tilde{\mathcal{H}}_1(d)$ can be rewritten as

$$h_1(\tilde{\mathbf{x}}; \tilde{\mathbf{w}}, b) = \sum_{i=1}^{d+1} \sum_{j=1}^{d+1} \psi_{ij}(\tilde{\mathbf{w}}, b) \phi_{ij}(\tilde{\mathbf{x}}) + \psi_0(\tilde{\mathbf{w}}, b) \phi_0(\tilde{\mathbf{x}}), \quad (24)$$

where

$$\begin{aligned}\psi_{ij} &= w_i w_j, \phi_{ij} = x_i x_j, 1 \leq i, j \leq d, \\ \psi_{d+1,j} &= b w_j, \psi_{i,d+1} = b w_i, \phi_{d+1,j} = x_j, \phi_{i,d+1} = x_i, 1 \leq i, j \leq d, \\ \psi_{d+1,d+1} &= b^2, \phi_{d+1,d+1} = 1, \psi_0 = -w_{d+1}^2, \phi_0 = \tilde{x}_{d+1}^2.\end{aligned}$$

Let $\phi(\tilde{\mathbf{x}}) = (\phi_{ij}(\tilde{\mathbf{x}}), \phi_0(\tilde{\mathbf{x}}))_{1 \leq i \leq d+1, 1 \leq j \leq d+1, i+j \leq 2d+2} \in \mathbb{R}^{(d+1)^2}$. Consider the linear class $\mathcal{L}_{(d+1)^2} = \{l(\mathbf{x}; \mathbf{w}, b) = \sum_{i=1}^{(d+1)^2} w_i x_i + b \mid (\mathbf{w}, b) \in \mathbb{R}^{(d+1)^2+1}\}$. Then for any $X_m = (\mathbf{x}_1, \dots, \mathbf{x}_m), y \in [-1, 1]$, the label patterns of $(\text{sgn}(h_1(\mathbf{x}_i) - y))_{i=1}^m$ that f_1 can achieve on X_m can also be achieved by $\mathcal{L}_{(d+1)^2}$ on $(\phi(\mathbf{x}_1), \dots, \phi(\mathbf{x}_m))$. Therefore, by definition we have

$$\begin{aligned}PDim(\mathcal{H}_1) &= VCdim(\{h(\mathbf{x}, y) = \text{sgn}(h_1(\mathbf{x}) - y) \mid h_1 \in \tilde{\mathcal{H}}_1(d), y \in [-1, 1]\}) \\ &\leq VCdim(\{h(\mathbf{x}, y) = \text{sgn}(l(\mathbf{x}) - y) \mid l \in \mathcal{L}_{(d+1)^2}, y \in [-1, 1]\}) \\ &= PDim(\mathcal{L}_{(d+1)^2}) \leq (d+1)^2 + 1,\end{aligned}\quad (25)$$

where inequality (25) holds by Theorem 11.6 (The Pseudo-Dimension of linear class) from [Anthony et al. \(1999\)](#).

Finally, from Lemma 2 we conclude that

$$\begin{aligned}PDim(\mathcal{H}_1 \otimes \mathcal{H}_2) &< 3(1 + \log PDim(\mathcal{H}_1) + \log PDim(\mathcal{H}_2))(PDim(\mathcal{H}_1) + PDim(\mathcal{H}_2)) \\ &< 3(1 + 3 \log(d+2))(d+2)(d+3) < 12(d+3)^2 \log(d+2),\end{aligned}$$

and therefore

$$\begin{aligned}
PDim(\mathcal{H}_{\mathcal{F}}) &\leq PDim(\mathcal{H}_1 \otimes \mathcal{H}_2 \otimes \mathcal{H}_3) \\
&< 3(1 + \log PDim(\mathcal{H}_1 \otimes \mathcal{H}_2) + \log PDim(\mathcal{H}_3))(PDim(\mathcal{H}_1 \otimes \mathcal{H}_2) + PDim(\mathcal{H}_3)) \\
&= 3(1 + \log(12(d+3)^2 \log(d+2)) + \log(d+1))(12(d+3)^2 \log(d+2) + d+1) \\
&< 3(1 + 6 \log(d+3))(13(d+3)^2 \log(d+3)) \\
&< 273(d+3)^2 \log^2(d+3).
\end{aligned}$$

Case-2: When $\mathcal{F}$ is a piece-wise linear function class with each instance constitutes m linear functions, i.e.,

$$\mathcal{F} = \{f(\mathbf{x}) = \mathbb{I}[\mathbf{w}_1^\top \mathbf{x} + b_1 \leq 0] \vee \dots \vee \mathbb{I}[\mathbf{w}_m^\top \mathbf{x} + b_m \leq 0] | (\mathbf{w}_i, b_i) \in \mathbb{R}^{d+1}, 1 \leq i \leq m\},$$

We first upper bound the VC-dimension of $\mathcal{F}$. If we take $\mathcal{F}, \mathcal{F}'$ to be binary function classes in (32) from Lemma 2, the Pdim of $\mathcal{F}$ coincides with the VCDim of $\mathcal{F}$. Hence, for the composition of m linear functions with each VCDim bounded by $d+1$, the VCDim of the new function is bounded by

$$\tilde{\mathcal{O}}(3(3(3(d+d)+d)+d)+\dots) \leq \tilde{\mathcal{O}}(3^m \cdot md) \leq \tilde{\mathcal{O}}(d \cdot 3^m),$$

where $\tilde{\mathcal{O}}$ denotes the big $\mathcal{O}$ notation omitting the log terms.

By Definition 1 and 2, we can compute each function $h \in \mathcal{H}_{\mathcal{F}}$ as

$$\begin{aligned}
h(f, \mathbf{e}; \mathbf{x}, c) &= D(\perp; (\mathbf{x}, c), \mathbf{e}) - D(f; (\mathbf{x}, c), \mathbf{e}) \\
&= \prod_{i=1}^m \mathbb{I}[\mathbf{w}_i^\top \mathbf{x}_i \leq 0] \cdot \left(\left\| \frac{\mathbf{e}}{2c} \right\|_2^2 - \min \left\{ \min_{1 \leq i \leq m} \left\{ \mathbb{I} \left[\mathbf{w}_i^\top \mathbf{x} + b_i > -\frac{\mathbf{w}_i^\top \mathbf{e}}{2c} \right] \cdot \left\| \mathcal{P}^{(i)} \left(\mathbf{x} + \frac{\mathbf{e}}{2c} \right) - \mathbf{x} \right\|_2^2 \right\}, \right. \right. \\
&\quad \left. \min_{1 \leq i, j \leq m} \left\{ \mathbb{I} \left[\mathbf{w}_i^\top \mathbf{x} + b_i > -\frac{\mathbf{w}_i^\top \mathbf{e}}{2c} \right] \cdot \mathbb{I} \left[\mathbf{w}_j^\top \mathbf{x} + b_j > -\frac{\mathbf{w}_j^\top \mathbf{e}}{2c} \right] \cdot \left\| \mathcal{P}^{(i,j)} \left(\mathbf{x} + \frac{\mathbf{e}}{2c} \right) - \mathbf{x} \right\|_2^2 \right\}, \right. \\
&\quad \left. \left. \min_{1 \leq i, j, k \leq m} \left\{ \prod_{t \in \{i, j, k\}} \mathbb{I} \left[\mathbf{w}_t^\top \mathbf{x} + b_t > -\frac{\mathbf{w}_t^\top \mathbf{e}}{2c} \right] \cdot \left\| \mathcal{P}^{(i,j,k)} \left(\mathbf{x} + \frac{\mathbf{e}}{2c} \right) - \mathbf{x} \right\|_2^2 \right\}, \dots \right\} \right), \quad (26)
\end{aligned}$$

where operator $\mathcal{P}^{(i,j)}$ denotes the L_2 -projection onto the intersection of hyperplanes $l_i : \mathbf{w}_i^\top \mathbf{x} + b_i \leq 0$ and $l_j : \mathbf{w}_j^\top \mathbf{x} + b_j \leq 0$, and $\mathcal{P}^{(i,j,k)}$ denotes the L_2 -projection onto the intersection of hyperplanes l_i, l_j, l_k , and so on. This is because there are in total 2^m possibilities in terms of the location of $\mathbf{x} + \frac{\mathbf{e}}{2c}$'s L_2 projection to the convex region denoted by $f(\mathbf{x}) = 1$, as $\mathcal{P}_f(\mathbf{x} + \frac{\mathbf{e}}{2c})$ can be on each hyperplane l_i , or on the intersections of any two l_i, l_j , or on the intersections of any three l_i, l_j, l_k , and so on.

Note that each $\mathcal{P}_f^{(r)}$ (i.e., the projection onto the intersection of r hyperplanes) has a closed-form which is a rational function with polynomial at most r . As a result, the Pseudo-dimension of the function class containing all functions like $\left\| \mathcal{P}^{(r)} \left(\mathbf{x} + \frac{\mathbf{e}}{2c} \right) - \mathbf{x} \right\|_2^2$ is at most $\mathcal{O}((rd)^r)$, since rd is the number of parameters each function has. Apply Eq. (31) in Lemma 2, we know the Pdim of the class $\mathbb{I} \left[\mathbf{w}_t^\top \mathbf{x} + b_t > -\frac{\mathbf{w}_t^\top \mathbf{e}}{2c} \right] \cdot \left\| \mathcal{P}^{(r)} \left(\mathbf{x} + \frac{\mathbf{e}}{2c} \right) - \mathbf{x} \right\|_2^2$ is at most $\tilde{\mathcal{O}}(3(d + (rd)^r))$. Continue to apply Eq. (32), we can upper bound the min of at most C_m^r functions with a Pdim of each at most $\tilde{\mathcal{O}}(3(d + (rd)^r))$ as $3^{C_m^r} \cdot C_m^r \cdot \tilde{\mathcal{O}}(3(d + (rd)^r))$, and the Pdim upper bound for the min of $(m+1)$ functions with a Pdim of

each at most $3^{C_m^r} \cdot C_m^r \cdot \tilde{O}(3(d + (rd)^r)), 1 \leq r \leq m$ is

$$\begin{aligned} Pdim(\mathcal{H}_{\mathcal{F}}) &\leq \mathcal{O}(d \cdot 3^m) \cdot \tilde{O}\left(3^m \sum_{r=0}^m 3^{C_m^r} \cdot C_m^r \cdot \tilde{O}(3(d + (rd)^r))\right) \\ &\leq \tilde{O}(d \cdot 3^m) \cdot \tilde{O}(3^{2^m}) \cdot \tilde{O}((dm)^m) \leq \tilde{O}(d^{m+1} \cdot 3^{2^m}). \end{aligned}$$

Case-3: When $\mathcal{F} = \{f(\mathbf{x}) = \mathbb{I}[\mathbf{w}^\top \phi(\mathbf{x}) + b \leq 0] | (\mathbf{w}, b) \in \mathbb{R}^{d+1}\}$ is the linear functions class with some feature transformation mapping ϕ , the best response of $(\mathbf{x}, c)$ is the solution of the following OP

$$\begin{aligned} \text{find} \quad & \mathbf{z}^* = \arg \min_{\mathbf{z}} \left\{ \left\| \mathbf{z} - \left(\mathbf{x} + \frac{\mathbf{e}}{2c} \right) \right\|_2^2 \right\} \\ \text{subject to} \quad & \mathbf{w}^\top \phi(\mathbf{z}) + b \leq 0. \end{aligned} \quad (28)$$

Since ϕ is invertible, it is equivalent to

$$\begin{aligned} \text{find} \quad & \mathbf{z}^* = \arg \min_{\mathbf{z}} \left\{ \left\| \phi^{-1}(\mathbf{y}) - \phi^{-1}\left(\phi\left(\mathbf{x} + \frac{\mathbf{e}}{2c}\right)\right) \right\|_2^2 \right\} \\ \text{subject to} \quad & \mathbf{w}^\top \mathbf{y} + b \leq 0. \end{aligned} \quad (29)$$

And also because ϕ preserves the order of pair-wise L_2 distance of any set of points, the solution of OP (29) is equivalent to the solution of

$$\begin{aligned} \text{find} \quad & \mathbf{z}^* = \arg \min_{\mathbf{z}} \left\{ \left\| \mathbf{y} - \phi\left(\mathbf{x} + \frac{\mathbf{e}}{2c}\right) \right\|_2^2 \right\} \\ \text{subject to} \quad & \mathbf{w}^\top \mathbf{y} + b \leq 0. \end{aligned} \quad (30)$$

As a result, we can compute $\mathbf{z}^*$ the same way as in Case-1 and the VCDim, PDim upper bounds in Case-1 still applies.

□

D.2 LEMMAS USED IN THE PROOF OF PROPOSITION 3 AND THEIR PROOFS

Lemma 2. For any class of real valued functions $\mathcal{F}, \mathcal{F}' \subseteq \{f : \mathbb{R}^d \rightarrow [-1, 1]\}$ and binary valued functions $\mathcal{G} \subseteq \{g : \mathbb{R}^d \rightarrow \{0, 1\}\}$, define $\mathcal{F} \otimes \mathcal{G} = \{h(\mathbf{x}) = f(\mathbf{x}) \times g(\mathbf{x}) | f \in \mathcal{F}, g \in \mathcal{G}\}$, and $\mathcal{F} \ominus \mathcal{F}' = \{h(\mathbf{x}) = \min\{f(\mathbf{x}), f'(\mathbf{x})\} | f \in \mathcal{F}, f' \in \mathcal{F}'\}$. Then, it holds that

$$PDim(\mathcal{F} \otimes \mathcal{G}) < 3(1 + \log d_{\mathcal{F}} d_{\mathcal{G}})(d_{\mathcal{F}} + d_{\mathcal{G}}), \quad (31)$$

$$PDim(\mathcal{F} \ominus \mathcal{F}') < 3(1 + \log d_{\mathcal{F}} d_{\mathcal{F}'})(d_{\mathcal{F}} + d_{\mathcal{F}'}), \quad (32)$$

where $d_{\mathcal{F}} = PDim(\mathcal{F}), d_{\mathcal{F}'} = PDim(\mathcal{F}'), d_{\mathcal{G}} = PDim(\mathcal{G})$.

Proof. By definition, $PDim(\mathcal{F})$ can be reduced to the VC dimension of the epigraph of $\mathcal{F}$, i.e.,

$$PDim(\mathcal{F}) = VCdim(\{h(\mathbf{x}, y) = \text{sgn}(f(\mathbf{x}) - y) | f \in \mathcal{F}, y \in [-1, 1]\}). \quad (33)$$

Let $\mathcal{X} = \mathbb{R}^d \times [-1, 1]$, consider an arbitrary set of points $X_m = \{(\mathbf{x}_i, y_i) \in \mathcal{X}\}_{i=1}^m$ with cardinality m and any binary hypothesis class $\mathcal{H} \subseteq \{h : \mathcal{X} \rightarrow \{0, 1\}\}$. Define the maximum shattering number

$$\Pi(m, \mathcal{H}) = \max_{X_m \in \mathcal{X}^m} \{\text{Card}\{(h(\mathbf{x}_1, y_1), \dots, h(\mathbf{x}_m, y_m)) \in \{0, 1\}^m | h \in \mathcal{H}\}\}$$

as the total number of label patterns that $\mathcal{H}$ can possibly achieve on $\mathcal{X}$. Next we upper bound the $\Pi(m, \{f * g | f \in \mathcal{F}, g \in \mathcal{G}\})$. For any fixed $X_m = \{(\mathbf{x}_i, y_i) \in \mathcal{X}\}_{i=1}^m \in \mathcal{X}^m$, we claim that the binary variable $\text{sgn}(f(\mathbf{x}_i)g(\mathbf{x}_i) - y_i)$ is determined by three binary variables $\text{sgn}(f(\mathbf{x}_i) - y_i)$ and $g(\mathbf{x}_i)$. This is because:

1. when $y_i \geq 0$, $f(\mathbf{x}_i)g(\mathbf{x}_i) \geq y_i$ holds if and only if $f(\mathbf{x}_i) \geq y_i$ and $g(\mathbf{x}_i) = 1$.
2. when $y_i < 0$, $f(\mathbf{x}_i)g(\mathbf{x}_i) \geq y_i$ holds if and only if $f(\mathbf{x}_i) \geq y_i$ and $g(\mathbf{x}_i) = 1$, or $g(\mathbf{x}_i) = 0$.

Therefore, any possible label pattern $(\text{sgn}(f(\mathbf{x}_1)g(\mathbf{x}_1) - y_1), \dots, \text{sgn}(f(\mathbf{x}_m)g(\mathbf{x}_m) - y_m)) \in \{0, 1\}^m$ is completely determined by the label patterns $(\text{sgn}(f(\mathbf{x}_1) - y_1), \dots, \text{sgn}(f(\mathbf{x}_m) - y_m))$ and $(g(\mathbf{x}_1), \dots, g(\mathbf{x}_m))$. As a result, it holds that

$$\begin{aligned} & \text{Card}\{(\text{sgn}(f(\mathbf{x}_1)g(\mathbf{x}_1) - y_1), \dots, \text{sgn}(f(\mathbf{x}_m)g(\mathbf{x}_m) - y_m)) | f \in \mathcal{F}, g \in \mathcal{G}\} \\ & \leq \text{Card}\{(\text{sgn}(f(\mathbf{x}_1) - y_1), \dots, \text{sgn}(f(\mathbf{x}_m) - y_m)) | f \in \mathcal{F}\} \times \text{Card}\{(g(\mathbf{x}_1), \dots, g(\mathbf{x}_m)) | g \in \mathcal{G}\}, \end{aligned}$$

which implies

$$\Pi(m, \mathcal{F} \otimes \mathcal{G}) \leq \Pi(m, \mathcal{F}) \times \Pi(m, \mathcal{G}). \quad (34)$$

Using the same argument, we can similarly show that

$$\Pi(m, \mathcal{F} \ominus \mathcal{F}') \leq \Pi(m, \mathcal{F}) \times \Pi(m, \mathcal{F}'). \quad (35)$$

Therefore, to show Eq. (31) and (32), it suffices to show Eq. (31) starting from Eq. (34). According to Sauer–Shelah Lemma (Sauer, 1972; Shelah, 1972), we have

$$\Pi(m, \mathcal{F}) \leq \sum_{i=0}^{VC(\mathcal{F})} \binom{m}{i} \leq \max\{m+1, m^{d_{\mathcal{F}}}\}, \quad (36)$$

where $VC(\mathcal{F})$ denotes the VC dimension of class $\{\text{sgn}(f(x) - y) | f \in \mathcal{F}\}$, which is also the Pseudo dimension of $\mathcal{F}$ (i.e., $d_{\mathcal{F}}$). And the second inequality of Eq. (36) holds because

1. when $d \geq 3$, we have

$$\left(\frac{d}{m}\right)^d \sum_{i=0}^d \binom{m}{i} \leq \sum_{i=0}^d \left(\frac{d}{m}\right)^i \binom{m}{i} \leq \sum_{i=0}^m \left(\frac{d}{m}\right)^i \binom{m}{i} = \left(1 + \frac{d}{m}\right)^m \leq e^d,$$

$$\text{and therefore } \sum_{i=0}^d \binom{m}{i} \leq \left(\frac{em}{d}\right)^d < m^d.$$

2. when $d = 2$, we have

$$\sum_{i=0}^d \binom{m}{i} = 1 + m + \frac{m(m-1)}{2} \leq m^2, \forall m \geq 2.$$

3. when $d = 1$, we have $\sum_{i=0}^d \binom{m}{i} = 1 + m$.

From Eq. (34) we know $\mathcal{F} \otimes \mathcal{G}$ has bounded Pseudo dimension. Suppose $PDim(\mathcal{F} \otimes \mathcal{G}) = d$, then by definition, there exists a set $\mathcal{Y}$ with cardinality d such that $\Pi(d, \mathcal{F} \otimes \mathcal{G}) = 2^d$. Therefore, from Eq (36) and (34) we have when $d \geq 2$,

$$2^d = \Pi(d, \mathcal{F} \otimes \mathcal{G}) \leq \Pi(d, \mathcal{F}) \times \Pi(d, \mathcal{G}) \leq \max\{d+1, d^{d_{\mathcal{F}}}\} \cdot \max\{d+1, d^{d_{\mathcal{G}}}\}. \quad (37)$$

For simplicity of notations we denote $d_1 = d_{\mathcal{F}}$, $d_2 = d_{\mathcal{G}}$ and without loss of generality assume $d_1 \geq d_2$. To complete our proof we need the following auxiliary technical Lemma 3, whose proof can be found in Appendix.

Lemma 3. For any $a \geq 2$ and $m > \frac{1.59a}{\ln 2} (\ln a - \ln \ln 2)$, it holds that $2^m > m^a$.

Now we are ready to prove our claim. Consider the following situations:

1. if $d_1 \geq 2, d_2 \geq 2$, from Eq (37) we obtain

$$2^d \leq d^{d_1+d_2}.$$

However, from Lemma 3 we know that when $d > \frac{1.59}{\ln 2}(d_1+d_2)(\ln(d_1+d_2)-\ln \ln 2)$, $2^d > d^{d_1+d_2}$ always holds. Hence, in this case we conclude $d \leq \frac{1.59}{\ln 2}(d_1+d_2)(\ln(d_1+d_2)-\ln \ln 2) < 2.3(d_1+d_2)(\ln(d_1+d_2)+0.37)$.

2. if $d_1 \geq 2, d_2 = 1$, from Eq (37) we obtain

$$2^d \leq (d+1)d^{d_1} < d^{d_1+2}.$$

From Lemma 3 we know that when $d > \frac{1.59}{\ln 2}(d_1+2)(\ln(d_1+2)-\ln \ln 2)$, $2^d > d^{d_1+2}$ always holds. Hence, in this case we conclude $d \leq \frac{1.59}{\ln 2}(d_1+2)(\ln(d_1+2)-\ln \ln 2) < 2.3(d_1+2)(\ln(d_1+2)+0.37)$.

3. if $d_1 = d_2 = 1$, from Eq (37) we obtain

$$2^d \leq (d+1)^2.$$

Since $2^m > (m+1)^2$ holds for any $m \geq 6$, we conclude that $d \leq 5 < \frac{1.59}{\ln 2}(2+2)(\ln(2+2)-\ln \ln 2)$.

Combining the three cases, we conclude that

$$\begin{aligned} PDim(\mathcal{F} \otimes \mathcal{G}) &< 2.3(\max\{2, d_1\} + \max\{2, d_2\})(\log(\max\{2, d_1\} + \max\{2, d_2\}) + 0.37) \\ &< 3(1 + \log d_1 d_2)(d_1 + d_2) \end{aligned}$$

□

Now we are ready to prove Theorem 1.

Proof of Theorem 1. Classic results from learning theory Pollard (1984) show the following generalization guarantees: Suppose $[0, H]$ is the range of functions in hypothesis class $\mathcal{H}$. For any $\delta \in (0, 1)$, and any distribution $\mathcal{D}$ over $\mathcal{X}$, with probability $1 - \delta$ over the draw of $S \sim \mathcal{D}^n$, for all functions $h \in \mathcal{H}$, the difference between the average value of h over S and its expected value gets bounded as follows:

$$\left| \frac{1}{n} \sum_{x \in S} h(x) - \mathbf{E}_{y \sim \mathcal{D}}[h(y)] \right| = \mathcal{O} \left(H \sqrt{\frac{1}{n} \left(PDim(\mathcal{H}) + \ln \left(\frac{1}{\delta} \right) \right)} \right)$$

Substituting $\mathcal{H}$ with the class of social distortion mitigation functions $\mathcal{H}_{\mathcal{F}} = \{h(f; \mathbf{x}, c) | f \in \mathcal{F}\}$ induced by some moderator function class $\mathcal{F} = f$ gives:

$$\left| \frac{1}{n} \sum_{(\mathbf{x}_i, c_i) \in S} h(f; \mathbf{x}_i, c_i) - \mathbb{E}_{(\mathbf{x}, c) \sim \mathcal{X} \times \mathcal{C}}[h(f; \mathbf{x}, c)] \right| = \mathcal{O} \left(H \sqrt{\frac{1}{n} \left(PDim(\mathcal{H}_{\mathcal{F}}) + \ln \left(\frac{1}{\delta} \right) \right)} \right)$$

Therefore, for a training set S of size $\mathcal{O} \left(\frac{H^2}{\varepsilon^2} [PDim(\mathcal{H}_{\mathcal{F}}) + \ln(1/\delta)] \right)$, the empirical average social distortion and the average social distortion on the distribution are within an additive factor of ε .

Next, we show for any class $\mathcal{F}$, distribution $\mathcal{D}$ over $\mathcal{X} \times \mathcal{C}$, if a large enough training set S is drawn from $\mathcal{D}$, then with high probability, every $f \in \mathcal{F}$, filters out approximately the same fraction of examples from the training set and the underlying distribution $\mathcal{D}$ had these examples manipulated to their ideal location(z'). In order to prove this, we use uniform convergence guarantees.

Let $\mathcal{X}' = \{\mathbf{x} + \frac{\mathbf{e}}{2c} \mid \mathbf{x} \in \mathcal{X}, c \in \mathcal{C}\}$ and $\mathcal{Y} = \{0\}$ and consider the joint distribution $\mathcal{D}'$ on $\mathcal{X}' \times \mathcal{Y}$. A hypothesis $f \in \mathcal{F}$ incurs a mistake on an example $(\mathbf{x} + \frac{\mathbf{e}}{2c}, y)$ if it labels it as positive or equivalently if it filters it out. By uniform convergence guarantees, given a training sample S' of size $\mathcal{O}(\frac{1}{\varepsilon^2}[VCDim(\mathcal{F}) + \log(1/\delta)])$, with probability at least $1 - \delta$ for every $f \in \mathcal{F}$, $|\text{err}_{\mathcal{D}'}(f) - \text{err}_S(f)| \leq \varepsilon$. This is equivalent to saying the fraction of points filtered out by f from $\mathcal{D}'$ and S' are within an additive factor of ε .

Here, $\mathcal{F}$ is the class of moderator functions, and $\mathcal{H}_{\mathcal{F}}$ is the class of social distortion mitigation functions induced by $\mathcal{F}$. Now, given a training set S of size $\mathcal{O}(\frac{1}{\varepsilon^2}[H^2(\text{PDim}(\mathcal{H}_{\mathcal{F}}) + \ln(1/\delta)) + VCDim(\mathcal{F})])$, by an application of union bound, for every $f \in \mathcal{F}$, the probability that the average social distortion of f on S and $\mathcal{D}$ differ by more than ε or the fraction of filtered points differ by more than ε is at most 2δ . This completes the proof. $\square$

E OMITTED PROOFS IN SECTION 5

E.1 PROOF OF LEMMA 1

Proof. Plugin the expression of $\mathbf{x}^* = \Delta(\mathbf{x}, c; \mathbf{e}, f)$ given by Proposition 1 and note that $\Delta(\mathbf{x}_i, c_i; \mathbf{e}, \perp) = \mathbf{x}_i + \frac{\mathbf{e}}{2c_i}$, we get a closed form of $DM(f; \mathcal{X})$ as shown below:

$$\begin{aligned} DM((\mathbf{w}, b); \mathcal{X}) &= \sum_{i=1}^n \{D(\perp; (\mathbf{x}_i, c_i), \mathbf{e}) - D(\mathbf{w}, b; (\mathbf{x}_i, c_i), \mathbf{e})\} \\ &= \sum_{i \in I_0(\mathbf{w}, b)} \left\| \frac{\mathbf{e}}{2c_i} \right\|_2^2 - \sum_{i \in I_1(\mathbf{w}, b)} \left\| \frac{\mathbf{e}}{2c_i} \right\|_2^2 - \sum_{i \in I_2(\mathbf{w}, b)} \left\| \frac{\mathbf{e}}{2c_i} - \frac{\mathbf{w}^\top (\mathbf{e} + 2c_i \mathbf{x}_i) \mathbf{w}}{2c_i \mathbf{w}^\top \mathbf{w}} - \frac{b \mathbf{w}}{\mathbf{w}^\top \mathbf{w}} \right\|_2^2 \\ &= \frac{1}{4} \sum_{i \in I_0} \frac{1}{c_i^2} - \frac{1}{4} \sum_{i \in I_1} \frac{1}{c_i^2} - \sum_{i \in I_2} \left\{ \frac{1}{4c_i^2} + \frac{1}{\mathbf{w}^\top \mathbf{w}} \left[(\mathbf{w}^\top \mathbf{x}_i + b)^2 - \left(\frac{\mathbf{w}^\top \mathbf{e}}{2c_i} \right)^2 \right] \right\} \\ &= \sum_{i \in I_2} \frac{1}{\mathbf{w}^\top \mathbf{w}} \left[-(\mathbf{w}^\top \mathbf{x}_i + b)^2 + \left(\frac{\mathbf{w}^\top \mathbf{e}}{2c_i} \right)^2 \right]. \end{aligned} \quad (38)$$

Here the set $I_0 = \{i \in [n] : \mathbf{w}^\top \mathbf{x}_i + b \leq 0\}$ contains the indices of all users who are marked as non-problematic and $I_1 = \{i \in [n] : \mathbf{w}^\top \cdot \left(\mathbf{x}_i + \frac{\mathbf{e}}{2c_i} \right) + b \leq 0\} \cap I_0$, $I_2 = \{i \in [n] : \mathbf{w}^\top \cdot \left(\mathbf{x}_i + \frac{\mathbf{e}}{2c_i} \right) + b > 0\} \cap I_0$.

Since a re-scaling of the vector $(\mathbf{w}, b)$ does not change the value of the RHS of Eq. (38), we may without loss of generality assume $\|\mathbf{w}\|_2 = 1$ and the DM function becomes

$$DM((\mathbf{w}, b); \mathcal{X}) = \sum_{i \in I_2} \left[-(\mathbf{w}^\top \mathbf{x}_i + b)^2 + \left(\frac{\mathbf{w}^\top \mathbf{e}}{2c_i} \right)^2 \right]. \quad (39)$$

Next, we argue that maximizing Eq. (39) under the constraint $\|\mathbf{w}\|_2 = 1$ is equivalent to maximizing it under the constraint $\|\mathbf{w}\|_\infty = 1$. This is because, for any solution $(\mathbf{w}^*, b^*)$ that yields the optimal value of

Eq. (39) with $\|\mathbf{w}^*\|_2 = 1$, re-scaling $(t\mathbf{w}^*, tb^*)$ such that $\|t\mathbf{w}^*\|_\infty = 1$ would also yield the largest value of the RHS of Eq. (38). And on the other hand, any solution $(\mathbf{w}^*, b^*)$ that yields the optimal value of Eq. (39) with $\|\mathbf{w}^*\|_\infty = 1$, we can also re-scale it such that $\|\mathbf{w}^*\|_2 = 1$. This suggests that we can equivalently consider the objective function given in Eq. (39) and replacing the original constraint $\|\mathbf{w}^*\|_2 = 1$ with $\|\mathbf{w}^*\|_\infty = 1$.

□

E.2 PROOF OF THEOREM 2

The ϵ -relaxed version of OP (9) is given by the following:

$$\begin{aligned} \text{find} \quad & \arg \min_{\mathbf{w}, b} \left\{ \sum_{i: \epsilon - \frac{\mathbf{w}^\top \mathbf{e}}{2c_i} \leq \mathbf{w}^\top \mathbf{x}_i + b \leq 0} \left[(\mathbf{w}^\top \mathbf{x}_i + b)^2 - \left(\frac{\mathbf{w}^\top \mathbf{e}}{2c_i} \right)^2 \right] \right\} \\ \text{subject to} \quad & \sum_{i=1}^n \mathbb{I} \left[\mathbf{w}^\top \left(\mathbf{x}_i + \frac{\mathbf{e}}{2c_i} \right) + b \leq 0 \right] \geq n - K, \\ & -1 \leq w_j \leq 1, 1 \leq j \leq d. \end{aligned} \quad (40)$$

Proof. For arbitrary n points $\mathbf{y}_1, \dots, \mathbf{y}_n \in \mathbb{R}^d$ and $K < n$, construct an OP (40) instance by letting $\mathbf{e} = (0, \dots, 0, 1)$, $c_1 = \dots = c_{2n} = \frac{1}{2\epsilon} > 0$, and

$$\mathbf{x}_i = \begin{cases} (y_{i,1}, \dots, y_{i,d}, 0), & 1 \leq i \leq n, \\ (0, \dots, 0, 0), & n+1 \leq i \leq 2n. \end{cases}$$

Then solving OP (40) is equivalent to

$$\begin{aligned} \text{find} \quad & \arg \max_{\mathbf{w}, b} \left\{ \sum_{\{1 \leq i \leq 2n: \epsilon(1 - \mathbf{w}^\top \mathbf{e}) \leq \mathbf{w}^\top \mathbf{x}_i + b \leq 0\}} \left[-(\mathbf{w}^\top \mathbf{x}_i + b)^2 + \epsilon^2 (\mathbf{w}^\top \mathbf{e})^2 \right] \right\} \\ \text{subject to} \quad & \sum_{i=1}^{2n} \mathbb{I} \left[\mathbf{w}^\top (\mathbf{x}_i) + b \leq -\epsilon \mathbf{w}^\top \mathbf{e} \right] \geq 2n - K, \\ & -1 \leq w_j \leq 1, \forall 1 \leq j \leq d+1, \\ & \mathbf{w} \neq \mathbf{0}. \end{aligned} \quad (41)$$

We argue that that the optimal $\mathbf{w}^*$ for OP (41) must satisfy $w_{d+1}^* = 1$, because for any $\mathbf{w}$ with $w_{d+1} < 1$, we have $\epsilon(1 - \mathbf{w}^\top \mathbf{e}) > 0$ and thus the set $\{1 \leq i \leq 2n : \epsilon(1 - \mathbf{w}^\top \mathbf{e}) \leq \mathbf{w}^\top \mathbf{x}_i + b \leq 0\}$ is empty. This means that the objective value of OP (41) is zero. However, any $\mathbf{w}, b$ with $w_{d+1} = 1$ such that $\mathbf{w}^\top \mathbf{x}_i + b = 0$ for some i yields an objective value $\epsilon^2 > 0$. As a result, the optimal $\mathbf{w}^*$ that maximize the objective of OP (41) must satisfy $w_{d+1}^* = 1$. Therefore, we can without loss of generality let $\mathbf{w}^\top \mathbf{e} = 1$ and then solving OP (41) is equivalent to

$$\begin{aligned} \text{find} \quad & \arg \max_{\mathbf{w}, b} \left\{ \sum_{\{1 \leq i \leq 2n: \mathbf{w}^\top \mathbf{x}_i + b = 0\}} [\epsilon^2] \right\} \\ \text{subject to} \quad & \sum_{i=1}^{2n} \mathbb{I}[\mathbf{w}^\top \mathbf{x}_i + b \leq -\epsilon] \geq 2n - K, \\ & -1 \leq w_j \leq 1, \forall 1 \leq j \leq d+1, \\ & \mathbf{w} \neq \mathbf{0}. \end{aligned} \quad (42)$$

Let $\tilde{\mathbf{w}} = (w_1, \dots, w_d)$ be the first d dimensions of $\mathbf{w}$, then solving OP (42) is equivalent to solving the following

$$\begin{aligned} \text{find} \quad & \arg \max_{\tilde{\mathbf{w}}, b} \left\{ n \cdot \mathbb{I}[b = 0] + \sum_{\{1 \leq i \leq n\}} \mathbb{I}[\tilde{\mathbf{w}}^\top \mathbf{y}_i + b = 0] \right\} \\ \text{subject to} \quad & \sum_{i=1}^n \mathbb{I}[\tilde{\mathbf{w}}^\top \mathbf{y}_i + b \leq -\epsilon] \geq n - K, \\ & -1 \leq \tilde{w}_j \leq 1, \forall 1 \leq j \leq d, \\ & \tilde{\mathbf{w}} \neq \mathbf{0}. \end{aligned} \quad (43)$$

We argue that the optimal solution of OP (43) must satisfy $b^* = 0$. This is because when $b = 0$, any $\tilde{\mathbf{w}}$ that satisfies $\tilde{\mathbf{w}}^\top \mathbf{y}_i + b = 0$ for some i yields an objective value at least $n + 1$. However, if $b \neq 0$, any $\tilde{\mathbf{w}}$

in the feasible region would yield an objective value at most n . As a result, solving OP (43) is equivalent to solving the following

$$\begin{aligned} & \text{find} && \arg \max_{\mathbf{w}} \left\{ \sum_{\{1 \leq i \leq n\}} \mathbb{I}[\mathbf{w}^\top \mathbf{y}_i = 0] \right\} \\ & \text{subject to} && \sum_{i=1}^n \mathbb{I}[\mathbf{w}^\top \mathbf{y}_i \leq -\epsilon] \geq n - K, \\ & && -1 \leq w_j \leq 1, \forall 1 \leq j \leq d, \\ & && \mathbf{w} \neq \mathbf{0}. \end{aligned} \quad (44)$$

Next, we show that optimizing Equation (44) is an NP-hard problem by showing the following decision problem that we call *ϵ -maximum feasible linear subsystem (MAX-FLS) with mandatory constraints* is NP-hard. Given $\epsilon > 0$ and a system of linear equations, with a mandatory set of constraints $A\mathbf{x} = \mathbf{b}_1$ where $\mathbf{b}_1 = [0, 0, \dots, 0]$ and an optional set of constraints $A\mathbf{x} \leq \mathbf{b}_2$ where $\mathbf{b}_2 = [-\epsilon, \dots, -\epsilon]$, and A is of size $d \times n$ and integers $1 \leq p \leq d$ and $0 \leq q \leq d$, does there exist a solution $\mathbf{x} \in \mathbb{R}^n$ satisfying at least p optional constraints while violating at most q mandatory constraints? Our proof is inspired by Amaldi & Kann (1995) that showed MAX-FLS is NP-hard, and we show that even when adding a set of mandatory constraints, it remains NP-hard.

In order to prove NP-hardness, we show a polynomial-time reduction from the known NP-complete *Exact 3-Sets Cover* that is defined as follows. Given a set S with $|S| = 3n$ elements and a collection $C = \{C_1, \dots, C_m\}$ of subsets $C_j \subseteq S$ with $|C_j| = 3$ for $1 \leq j \leq m$, does C contain an exact cover, i.e. $C' \subseteq C$ such that each element s_i of S belongs to exactly one element of C' ?

Let (S, C) be an arbitrary instance of *Exact 3-Sets Cover*. We will construct a particular instance of *ϵ -MAX-FLS with mandatory constraints* denoted by (A, b, p, q, ϵ) such that there exists an *Exact 3-Sets Cover* if and only if the answer to the *ϵ -MAX-FLS with mandatory constraints* instance is affirmative.

We construct an instance of *ϵ -MAX-FLS with mandatory constraints* as follows. There exists one variable x_j for each subset $C_j \in C$, $1 \leq j \leq m$. Equations (45) to (47) are optional and Equations (48) to (50) are mandatory constraints. Equation (45) are coverage constraints to make sure each element in S is covered. Constant $a_{i,j}$ is equal to 1 if $s_i \in C_j$ and is equal to 0 otherwise.

$$\sum_{j=1}^{|C|} a_{i,j} x_j - x_{m+1} = 0 \quad \forall 1 \leq i \leq 3n \quad (45)$$

$$x_j - x_{m+1} = 0 \quad \forall 1 \leq j \leq m \quad (46)$$

$$x_j = 0 \quad \forall 1 \leq j \leq m \quad (47)$$

$$\sum_{j=1}^{|C|} a_{i,j} x_j - x_{m+1} \leq -\epsilon \quad \forall 1 \leq i \leq 3n \quad (48)$$

$$x_j - x_{m+1} \leq -\epsilon \quad \forall 1 \leq j \leq m \quad (49)$$

$$x_j \leq -\epsilon \quad \forall 1 \leq j \leq m \quad (50)$$

We set $p = 3n + m$ and $q = 2n + 2m$. Now, in any solution $\mathbf{x}$, we must have $x_{m+1} \neq 0$, since $x_{m+1} = 0$ implies that $x_j = 0$ for all $1 \leq j \leq m$ in order to have at least $3n + m$ optional constraints satisfied. However, if all x variables are set as 0, then none of the mandatory constraints would be satisfied which is a contradiction.

Now, given any exact cover $C' \subseteq C$ of (S, C) , the vector $\mathbf{x}$ defined by:

$$x_j = \begin{cases} -\epsilon & \text{if } C_j \in C' \text{ or } j = m + 1 \\ 0 & \text{otherwise} \end{cases}$$

satisfies all equations of type Equation (45) and exactly m of Equations (46) and (47). Therefore, $\mathbf{x}$ satisfies $3n + m$ optional constraints in total. However, all constraints of type Equation (48) are violated. When $x_j = -\varepsilon$, the mandatory constraint $x_j \leq -\varepsilon$ is satisfied, and $x_j - x_{m+1} \leq -\varepsilon$ is violated. However, when $x_j = 0$, both constraints $x_j - x_{m+1} \leq -\varepsilon$, $x_j \leq -\varepsilon$ are violated. Since $|C'| = n$, the total number of mandatory constraints violated equals $3n + 2m - n = 2n + 2m$.

Conversely, suppose that we have a solution $\mathbf{x}$ that satisfies at least $3n + m$ optional constraints and violates at most $2n + 2m$ mandatory constraints. By construction, since $\mathbf{x}$ satisfies $3n + m$ optional constraints, it satisfies all constraints of type Equation (45) and exactly m constraints among Equations (46) and (47) (recall that we are interested in non-trivial solutions, therefore $x_{m+1} \neq 0$). This implies each x_j is either equal to x_{m+1} or 0. Now, consider the subset $C' \subseteq C$ defined by $C_j \in C'$ if and only if $x_j = x_{m+1}$. This gives an exact cover of (S, C) .

Finally, we can conclude that given solution $\mathbf{x}$ that satisfies at least $3n + m$ optional constraints and violates at most $2n + 2m$ mandatory constraints, the subset $C' \subseteq C$ defined by $C_j \in C'$ if and only if $x_j = x_{m+1}$ is an exact cover of (S, C) . $\square$

E.3 PROOF OF PROPOSITION 4

Proof. OP (10) is equivalent to the following problem

$$\text{find } \arg \min_{\mathbf{w}, b} \{-DM(\mathbf{w}, b) + \lambda P(\mathbf{w}, b)\}. \quad (51)$$

And according to the definitions we have

$$(\mathbf{w}_1, b_1) = \arg \min_{\mathbf{w}, b} \{-DM(\mathbf{w}, b) + \lambda_1 P(\mathbf{w}, b)\}, \quad (52)$$

$$(\mathbf{w}_2, b_2) = \arg \min_{\mathbf{w}, b} \{-DM(\mathbf{w}, b) + \lambda_2 P(\mathbf{w}, b)\}. \quad (53)$$

As a result,

$$\begin{aligned} -DM(\mathbf{w}_2, b_2) + \lambda_2 P(\mathbf{w}_2, b_2) &\leq -DM(\mathbf{w}_1, b_1) + \lambda_2 P(\mathbf{w}_1, b_1) \\ &= -DM(\mathbf{w}_1, b_1) + \lambda_1 P(\mathbf{w}_1, b_1) + (\lambda_2 - \lambda_1)P(\mathbf{w}_1, b_1) \\ &\leq -DM(\mathbf{w}_2, b_2) + \lambda_1 P(\mathbf{w}_2, b_2) + (\lambda_2 - \lambda_1)P(\mathbf{w}_1, b_1) \end{aligned} \quad (54)$$

Rearranging terms in Eq. (54) we have $(\lambda_2 - \lambda_1)(P(\mathbf{w}_2, b_2) - P(\mathbf{w}_1, b_1)) \leq 0$, which implies $P(\mathbf{w}_1, b_1) \leq P(\mathbf{w}_2, b_2)$.

Now we prove $DM(\mathbf{w}_1, b_1) \leq DM(\mathbf{w}_2, b_2)$. First of all, when $\lambda_2 = 0$, we have $(\mathbf{w}_2, b_2) = \arg \max_{\mathbf{w}, b} DM(\mathbf{w}, b)$ and therefore $DM(\mathbf{w}_1, b_1) \leq DM(\mathbf{w}_2, b_2)$ must hold. When $\lambda_2 > 0$, it also holds that

$$\begin{aligned} \frac{\lambda_1}{\lambda_2} (-DM(\mathbf{w}_2, b_2) + \lambda_2 P(\mathbf{w}_2, b_2)) &\leq \frac{\lambda_1}{\lambda_2} (-DM(\mathbf{w}_1, b_1) + \lambda_2 P(\mathbf{w}_1, b_1)) \\ &= -DM(\mathbf{w}_1, b_1) + \lambda_1 P(\mathbf{w}_1, b_1) - \frac{\lambda_1 - \lambda_2}{\lambda_2} DM(\mathbf{w}_1, b_1) \\ &\leq -DM(\mathbf{w}_2, b_2) + \lambda_1 P(\mathbf{w}_2, b_2) - \frac{\lambda_1 - \lambda_2}{\lambda_2} DM(\mathbf{w}_1, b_1) \end{aligned} \quad (55)$$

Rearranging terms in Eq. (55) we have $(\lambda_1 - \lambda_2)(DM(\mathbf{w}_1, b_1) - DM(\mathbf{w}_2, b_2)) \leq 0$, which implies $DM(\mathbf{w}_1, b_1) \leq DM(\mathbf{w}_2, b_2)$. $\square$

F ADDITIONAL EXPERIMENTS

F.1 ADDITIONAL DETAILS OF EXPERIMENT

Synthetic data generation: We generate synthetic dataset from mixed Gaussian distribution in $\mathbb{R}^d$ to mimic the distribution of $\mathbf{x}$. Specifically, we first sample k centers $\mathbf{c}_i$ from $\mathcal{N}(0, I_d)$ and then for each $\mathbf{c}_i$ generate $m = n/k$ samples from $\mathcal{N}(\mathbf{c}_i, \sigma_i^2 I_d)$, where σ_i is sampled uniformly at random from $[0.3, 0.5]$. Without loss of generality we set $\mathbf{e}$ as the unit vector $(1, 0, \dots, 0)$, and sample c_i independently from uniform distribution $\mathcal{U}[0.5, 1.5]$. In the experiments we choose $d = 5, n = 500, k = 5$, and additional result under different data scales can be found in Appendix F.

Optimization solver setup: we solve OP (11) by setting the ℓ_2 penalty $P(\mathbf{w}, b) = \|g(\mathbf{w}, b; U, \mathbf{e})\|_2$, i.e., $P_i(\mathbf{w}, b) = \mathbb{I}[y_i > 0] \cdot y_i^2$, where $a_i = \frac{\mathbf{w}^\top \mathbf{e}}{2c_i}, y_i = \mathbf{w}^\top \mathbf{x}_i + b + a_i$ as defined in the constraints of OP (11). The reason we choose such a form is because the resultant social good loss function l can preserve continuity and first-order differentiable property, paving the way for gradient-based method. Additional details about the optimizer setup can be found in Appendix F.1.

We solve OP (11) by setting the ℓ_2 penalty $P(\mathbf{w}, b) = \|g(\mathbf{w}, b; U, \mathbf{e})\|_2$, i.e., $P_i(\mathbf{w}, b) = \mathbb{I}[y_i > 0] \cdot y_i^2$, where $a_i = \frac{\mathbf{w}^\top \mathbf{e}}{2c_i}, y_i = \mathbf{w}^\top \mathbf{x}_i + b + a_i$ as defined in the constraints of OP (11). The reason we choose such a form is because the resultant social good loss function l can preserve continuity and first-order differentiable property, paving the way for gradient-based method. To further differentiate l at $y_i = 0$, we apply spline interpolation at $y_i = 0.1$ to round the corner at $y_i = 0$ while ensuring that $l \rightarrow 0$ as $y_i \rightarrow -\infty$. The surrogate loss function l for each user $(\mathbf{x}_i, c_i)$ after such regularizations compared with the true loss l is illustrated in the leftmost panel of Figure 3, and we can observe that the minimum of l is achieved when $y_i = a_i$, i.e., when the original feature $\mathbf{x}$ is on the decision boundary of filter f . In addition, for a larger penalty λ , moving across the boundary (i.e., y_i moving to the right side of a_i) would incur a larger and more rapidly increasing loss. The objective is thus the sum of these surrogate losses at all points $(\mathbf{x}_i, c_i)$. To account for the boundary constraint $-1 \leq w_i \leq 1$, we employ the standard projected gradient descent.

The surrogate loss $\tilde{l}(y, a; \epsilon)$ for a single point $(\mathbf{x}, c)$ we use in the experiment is given by the following explicit form:

$$\tilde{l}(y, a; \epsilon, \lambda) = \begin{cases} \frac{(1-\epsilon)^2 a^3}{2\epsilon y - 4a(1-\epsilon) + 3a(1-\epsilon)^2}, & y < (1-\epsilon)a, \\ y^2 - 2ay, & (1-\epsilon)a \leq y \leq a, \\ \lambda(y-a)^2 - a^2, & y > a, \end{cases} \quad (56)$$

where $y = \mathbf{w}^\top \mathbf{x} + b + a, a = \frac{\mathbf{w}^\top \mathbf{e}}{2c}$, as shown in Figure 4. In our experiments we choose $\epsilon = 0.9$ and use different λ ranging from 0.1 to 100.

Then we use projected gradient descent (PGD) to solve the following OP 57 with the exact gradient of $\tilde{l}$ w.r.t. $\mathbf{w}$ and b . The learning rate of PGD is set to 0.1 and the maximum iteration steps is set to 2000.

$$\begin{aligned} \text{find} \quad & \arg \min_{\mathbf{w}, b} \left\{ \sum_{1 \leq i \leq n} \tilde{l}(y_i, a_i) \right\} \\ \text{subject to} \quad & y_i = \mathbf{w}^\top \mathbf{x}_i + b + a_i, 1 \leq i \leq n, \\ & a_i = \frac{\mathbf{w}^\top \mathbf{e}}{2c_i}, 1 \leq i \leq n, \\ & -1 \leq w_i \leq 1, 1 \leq i \leq n. \end{aligned} \quad (57)$$

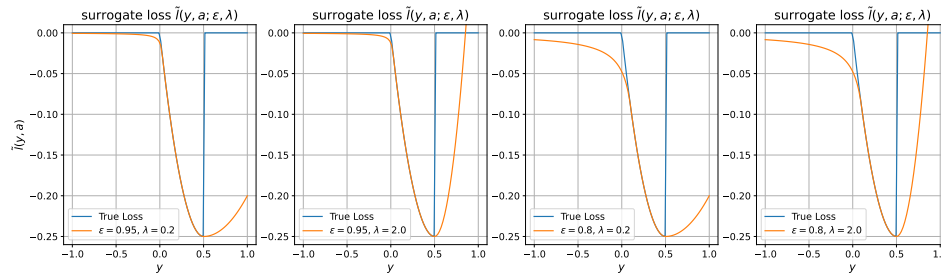


Figure 4: The constructed quasi-convex single point surrogate loss function with different smoothing parameter ϵ and soft freedom of speech penalty strength λ . In this illustration we set $a = 0.5$.

F.2 ADDITIONAL RESULT UNDER DIFFERENT DIMENSION d

We also plot the trade-offs achieved by the computed optimal linear moderators across different dimensions d , with the results shown in Figure 5. As the figure illustrates, higher dimensions introduce more noise into the results, but the same underlying insights remain observable.

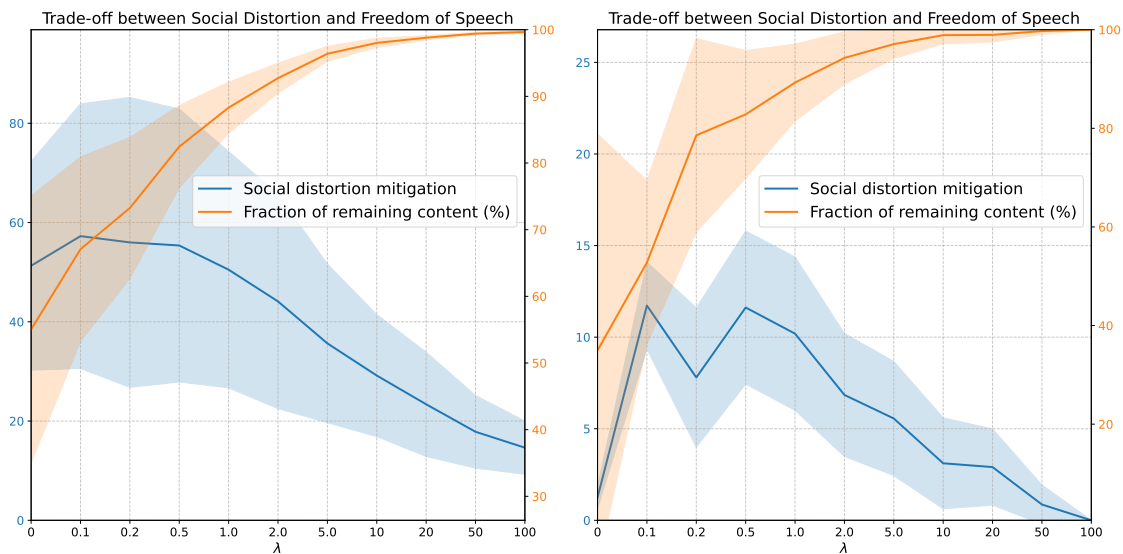


Figure 5: Social distortion mitigation (blue) and the fraction of remaining content on the platform (yellow) incurred by the computed moderator obtained under different $\lambda \in [0.1, 100]$. Left: $d = 2$, Right: $d = 10$. Error bars are 1σ region based on results from 20 independently generated datasets.