

The Order Beneath the Hostility

A Qualitative Study of the Moral Economy of Competitive Online Play

ANTON JANDE WALDAU
LINUS LÖNNQVIST

Bachelor Thesis
Stockholm school of Economics
2026



Abstract

Over 3.6 billion people play digital games, and the leading competitive games are free-to-play services whose revenue depends on player retention. Platforms have spent two decades trying to suppress hostile conduct between teammates and have not succeeded. Existing structural accounts explain why toxicity persists but cannot recover how community members themselves reason about it as legitimate; they read that reasoning as a cover story rather than as the substantive content of a moral order. This thesis asks: how do members of competitive online communities rationalize hostile conduct as normatively legitimate, and what does this reveal about the conditions under which platform governance fails to enter a community's moral order? An interpretivist, abductive design draws on twenty-two semi-structured interviews with active League of Legends ranked players. The framework synthesises Thompson's customs, Polanyi's commodity logic via Bolton and Laaser, and Sayer's lay morality into a three-layer moral economy adapted for symbolic-stakes, pseudonymous, transient communities. Players treat the skill rating as what earns them a voice in the match. They treat high skill as a licence to criticise teammates, and they apply that licence match by match by deciding who deserves blame. The community gives its highest standing to the high-skill player who holds the licence and chooses not to use it: the thesis names this figure apex restraint. Platform tools arrive in a community that already has rules for legitimate conduct, and players redirect each tool to their own ends. Women and players with origin-marked names carry the heaviest cost. Hostility is what the community's rules produce.

Keywords: *Competitive online gaming, hostile conduct, moral economy, customs, platform governance, player community.*

Authors

Anton Jande Waldau (25870)

Linus Lönnqvist (25942)

Supervisor

Jonathan Schunnesson, Affiliated Researcher, Department of Management and Organization

Examiner

Laurence Romani, Professor, Department of Management and Organization

Bachelor Thesis

Bachelor Program in Management

Stockholm School of Economics

© *Anton Jande Waldau and Linus Lönnqvist, 2026*

Acknowledgments

We thank the players who participated in our study. Without their time and energy, this thesis would not have been possible to write. They spoke openly about conduct the broader culture stigmatises, and the analysis that follows rests entirely on their candour. We also thank our supervisor Jonathan Schunnesson for his insightful discussions and guidance throughout the project, and our course director Laurence Romani for the pedagogical set-up of the course and the inspirational seminars that shaped how we approached the work. We thank our family and friends for their unconditional support.

Lastly, we thank the Stockholm School of Economics for its stimulating educational environment and the academic camaraderie.

Anton and Linus

Definitions

Twelve terms recurring throughout the thesis are defined below. Working definitions follow the cited source where applicable. Game-specific terms are defined as the empirical site uses them.

Term	Definition
Competitive online multiplayer game	A digital game in which two or more players, typically organised into teams, compete within a structured ruleset, often with persistent skill-rating systems and short-form match cycles.
MOBA	MOBA stands for <i>Multiplayer Online Battle Arena</i> , a subgenre of strategy games where two teams of players control unique characters and compete to destroy the opposing team's base.
League of Legends	Riot Games' team-based, free-to-play multiplayer online battle arena released in 2009; the empirical site of this thesis. Two teams of five players compete to destroy the opposing team's base.
Ranked mode / ranked ladder	The competitive game mode in which players are matched by skill rating, with each match updating a public ladder position. Distinguished from casual modes by the consequentiality of each match for one's standing.
Match	A single competitive game session, typically 25–45 minutes in the empirical site. Players are matched by skill rating and assigned to teams immediately before the match begins.
Mute	The in-game function that silences a specific player's typed chat for the rest of the match; the most common coping mechanism for hostile chat (Poeller et al., 2023).
Chat ban	A platform-imposed temporary restriction preventing a player from typing in match chat, applied automatically by Riot Games following sustained reports for hostile chat.
Toxicity / hostile conduct	Following Kordyaka et al. (2025), a multidimensional collective term for disruptive in-game conduct, varying across form (text, speech, behaviour), target (teammate, opponent), intention, and timing. Used interchangeably with hostile conduct throughout.
Griefing / griefer	Following Achterbosch et al. (2024), "a type of toxic behavior that focuses on player-to-player in-game disruption"; the perpetrator is termed a griefer, the targeted player the grieved. A

Term	Definition
	specific behavioural subset of toxicity, often operating through in-game mechanics rather than chat.
Tilt / tilting	A cyclical process in which a frustrating in-game event progressively impairs both performance and self-regulation, eventually producing toxic output (Türkay et al., 2020).
Moral economy	Following Thompson (1971, 1991) and Bolton and Laaser (2013), a community's shared system of legitimate expectations about how members should treat one another, made visible through outrage at violations and through the customary justifications offered for direct action.
Customs	In Thompson's (1991) technical sense, normative claims actively constructed and defended by a community, rather than paternalistic norms imposed from above.
Gender masking	A coping strategy in which female players conceal gender markers (voice, name, avatar) to avoid gender-targeted hostility (Fox & Tang, 2017).
Flagging	The player-facing post-match reporting function through which players can flag a teammate's or opponent's conduct as a rule breach for review by the platform's automated or human moderation system.
Patches	Scheduled game updates released by the developer, typically at intervals of two to four weeks, that adjust gameplay parameters, mechanics, and platform features. Patches are a primary vector through which platform governance and community customs change.
Slurflation	Modelled on monetary inflation, the loss of shock value that established slurs and hostile formulas suffer through overuse, desensitisation, and chat-filter restriction, leaving the community to escalate or invent new forms to land the same impact.
Slurfest	A collective pile-on in which several players direct escalating slurs at one target once a vulnerability marker, such as gender or national origin, becomes visible.

Table of contents

1. Introduction	8
1.1. The Economic Stakes	8
1.2. Purpose and Research Question	8
1.3. Scope	9
2. Literature Review	9
2.1. What is Toxicity?	9
2.2. What is an Online Gaming Community?	10
2.3. Disinhibition and Emotional Architecture	10
2.4. Toxic Meritocracy and Its Reproduction	10
2.5. The Gendered Distribution of Costs	11
2.6. Governance Failure	12
2.7. The Theoretical Gap and the Conversation This Thesis Joins	12
3. Theoretical Framework	13
3.1. From Thompson to Bolton and Laaser	13
3.2. The Three Layers Applied to Competitive Gaming	14
3.3. Adaptations	17
4. Methodology	17
4.1. Research Approach	18
4.2. Sampling and Recruitment	18
4.3. Interview Guide and Data Collection	19
4.4. Analytical Procedure	19
4.5. Research Quality and Ethics	22
4.6. Limitations	22
5. Empirical Findings	23
5.1. Worth in the Ladder: Skill and Rank	23
5.1.1. Skill as Licence	23
5.1.2. Rank as Identity and Identity Threat	23
5.1.3. Griefing as the one hard limit	23
5.2. Customs of In-Game Conduct	24
5.2.1. The In-Game Register	24
5.2.2. Reciprocity precondition	24
5.2.3. Creative cruelty as community currency	24
5.3. Lay Moral Reasoning in Practice	25
5.3.1. Targeting Direction by Game State	25
5.3.2. Ego-protective blame attribution	25
5.3.3. Cognitive depletion	25

5.3.4. Desensitisation	25
5.4. Toxicity as Strategy and the Asymmetry of Speech	26
5.4.1. Toxicity as Competitive Instrument	26
5.4.2. Game State and the Licence to Type	26
5.4.3. Performance-licensed prosocial speech	26
5.4.4. The Scepticism of Positivity	26
5.5. Platform Governance: External Compliance Without Internalisation	27
5.6. The Implicit Norm	27
5.6.1. Impunity Logic	27
5.6.2. Cost of marked national origin	27
5.6.3. Female visibility and concealment	27
6. Analysis	28
6.1. Skill-as-Licence as Moral Enforcement	28
6.2. Apex Restraint as the Order's Ideal	29
6.3. The Order Absorbs Governance	30
6.4. The Same Logic, Differently Distributed	31
7. Discussion and Conclusion	31
7.1. Answering the Research Question	31
7.1.1. Sustaining the Moral Order in Competitive Online Play	32
7.2. Contribution to the Academic Debate	32
7.3. Implications for Practitioners and Management	32
7.3.1. For platforms and game studios	32
7.3.2. For community managers and moderation system designers	32
7.3.3. For players outside the default	33
7.3.4. Ethical implications	33
7.4. Limitations	33
7.5. Future Research	33
7.6. Conclusion	33
References	35
Appendix 1. Interview Guide	38
A1.1. Opening and Framing	38
A1.2. Background and Entry into Competitive Play (Individual Layer)	38
A1.3. Match Conduct and Community Customs (Customs Layer)	39
A1.4. Platform Governance: Reporting, Peer-Endorsement System, Sanctions (Structural Layer)	40
A1.5. Reflection and Close	40
Appendix 2. Interview Specifications	41
Appendix 3. Use of AI Tools	42

1. Introduction

Every day, hundreds of millions of people log into competitive online games and play short matches with strangers. The matches are tense, the players are anonymous, and the atmosphere notoriously hostile: insults, slurs, and threats are common enough that the industry has its own word for them, *toxicity*. Game companies have spent two decades trying to suppress this conduct. They have deployed chat filters, reporting systems, automated sanctions, and rewards for good behaviour. The behaviour has not gone away. The communities the tools were built to discipline have absorbed the tools themselves: in-game chat filters become puzzles to evade, reporting functions become competitive weapons, peer-endorsement systems are repurposed to reward good *play* rather than good conduct. The community has turned the instruments of governance to its own ends and re-articulated the toxic conduct itself as legitimate. These platforms govern speech, conduct, and belonging for populations larger than most states. This thesis examines the process through which players legitimize that behaviour.

1.1. The Economic Stakes

The scale matters. Over 3.6 billion people play digital games as of 2025, generating \$188.8 billion in annual revenue (Newzoo, 2025; Icon Era, 2026), with the ten most popular online titles sustaining over 850 million monthly active players (Twinklfinite, 2025). Toxicity is structural here. Nearly half of gaming sessions involve hate or harassment, and toxic experiences are linked to stress, depression, anxiety, and social withdrawal (Zsila et al., 2022; ADL, 2025; Kordyaka et al., 2020; Wong & Ratan, 2024; Frommel et al., 2023). Women carry more of that cost than men: they face sexual harassment on top of general hostility, and many quit (Fox & Tang, 2017). The conduct is also a direct economic threat. Competitive games are free to play and sustained by voluntary microtransactions, with the average player spending \$147 annually (Icon Era, 2025); retention is everything. Retention is also asymmetric. Players who carry the cost of hostility leave. Players who impose it stay. The stakes here outgrow the firm: the same conduct that threatens financial revenue concerns the institutions tasked with global governance. Two of the United Nations' Sustainable Development Goals frame these stakes: SDG 5, on gender equality, and SDG 16, on institutional accountability for platforms governing speech and belonging. Governance failure here is not insufficient enforcement but enforcement that does not convert into compliance.

1.2. Purpose and Research Question

The purpose of this thesis is to recover the normative logic through which competitive online community members render hostile conduct legitimate, and to specify the conditions under which platform governance fails to enter that architecture. From this purpose we derive a single research question:

How do members of competitive online communities rationalize hostile conduct as normatively legitimate, and what does this reveal about the conditions under which platform governance fails to enter a community's moral order?

1.3. Scope

This study is bounded in three respects. First, it focuses on the game League of Legends. Toxicity is a feature of virtually all competitive online environments, but a single-site design buys depth that cross-platform comparisons would sacrifice. The findings travel to environments that share the structural conditions that make them relevant: anonymity, competitive pressure, low social presence. Second, participation was limited to players active in competitive play, since this environment concentrates performance pressure, accountability, and rank-related incentives. Including casual game modes would have introduced player behaviour shaped by different motivations and stakes, which fell outside the scope of the analysis. Third, the study examines player perception rather than behavioural or platform-level data. It makes no claims about the actual frequency of toxic incidents. What it investigates is how players construct meaning around this phenomenon, what they understand as legitimate, and how they reason about the governance tools they encounter.

2. Literature Review

Research on hostile conduct in competitive online multiplayer gaming has moved from cataloguing abusive acts to mapping the structural conditions that produce them. The field has established that toxic behaviour is widespread, that it inflicts measurable harm, that environmental conditions rather than deviant minorities produce it, and that the punitive moderation systems built to address it routinely fail. What it has not established is how community members themselves articulate the normative order through which hostile conduct becomes legitimate. The review proceeds in three steps. We define toxicity and the kind of community that produces it; we trace how that community's cultural logic reproduces and how its costs are unevenly borne; we locate the conversation about hostile conduct and platform governance to which this thesis contributes.

2.1. What is Toxicity?

Definitional precision has been a persistent problem. Reviewing 32 articles, Kordyaka, Karaosmanoglu, and Laato (2025) find that scholars use harassment, griefing, flaming, and toxicity inconsistently and overlap them, producing terminology that is “general and unspecific” (p. 2); they propose treating toxicity as a multidimensional collective term varying across form, target, intention, and timing. Within this vocabulary, *griefing* has its own tradition: Achterbosch, Vamplew, and March (2024) define it as “a type of toxic behavior that focuses on player-to-player in-game disruption,” distinguishing the perpetrator (the griever) from the

targeted player (the grieved). Griefing operates through gameplay rather than chat: deliberate underperformance, abandoning a match mid-play, refusing one's role, or exploiting game mechanics to harm teammates' experience. The framing distinguishes toxicity from cyberbullying, which is defined by repetition and a power imbalance among other criteria that competitive matchmaking eliminates (Kordyaka et al., 2025). Kou (2020) names five subtypes of toxicity in League of Legends: communicative aggression (most prominently flaming: sending hostile or insulting messages or using the report system as a coercive threat), cheating, hostage holding (keeping teammates "in an unpleasant situation", i.e. refusing to surrender a game), mediocritizing (playing to win, but in a way the team judges sub-optimal) and sabotaging. Several operate without language, exposing the limits of moderation systems that struggle to detect non-chat conduct.

2.2. What is an Online Gaming Community?

Online gaming environments are not simply digital extensions of offline communities. The structural conditions of online interaction (anonymity, invisibility, lack of simultaneous interaction, reduced eye contact) loosen the inhibitions that ordinarily constrain face-to-face conduct (Suler, 2004; Lapidot-Lefler & Barak, 2012). Liu and Agur (2023) document the practical effect: in interviews with players of a Chinese mobile team-based game, anonymity reduced perceived consequences ("after all they don't know me"), and toxicity fell when players used voice chat or played with familiar others (pp. 612–613). Kou (2021) characterises online games as a "synthetic world" bounded by specific norms and unique entry barriers (p. 334:5). Competitive ranked play in League of Legends, the empirical site of this thesis, pushes these conditions to their limit: matches last under an hour, teammates are pseudonymous (using fictitious names), and skill-based matchmaking forms (mostly) temporary teams (Poeller et al., 2023, p. 374:3), so the same teammate is rarely encountered twice. The community we examine is therefore pseudonymous and transient, reformed match by match.

2.3. Disinhibition and Emotional Architecture

Suler's (2004) account establishes that disinhibition is structurally available online; it does not explain why hostility is so consistently directed and defended. Türkay et al. (2020) and Poeller et al. (2023) supply the dynamic missing piece: *tilting*, in which a frustrating event progressively impairs self-regulation (Türkay et al., 2020, p. 5), and *desensitisation*, in which experienced players come to treat hostility as inherent to the form (Türkay et al., 2020, p. 5; Poeller et al., 2023, p. 374:5). Together these accounts reframe toxic conduct as an emergent product of the game's emotional architecture rather than as individual disposition.

2.4. Toxic Meritocracy and Its Reproduction

The cultural logic that organises gaming hostility is theorised by Paul, who coined "toxic meritocracy" to describe how online game cultures, facilitated by game design and practiced by

gamers, "tend to celebrate merit and put extreme focus on skill, achievement, and hard work" (Paul, 2018, as cited in Kou, 2021, p. 334:5). Where *ranked ladders* (skill-based ranking systems) and visible performance statistics constitute the symbolic order, unskilled players become legitimate objects of contempt; harassment of poor performers becomes an enactment of community norms rather than a deviation. Kou and Gui (2021) record the empirical signature: players treat underperformance and a "careless attitude" as flaggable in their own right, regardless of whether the player is also verbally toxic. Achterbosch et al. (2024) confirm the asymmetry across 656 online players: grieving increased perpetrators' autonomy and competence while reducing targeted players' autonomy and relatedness. Competitive contribution becomes a substitute for, and a measure of, conduct. Social cognitive theory illuminates how this logic reproduces. Sun et al. (2024), drawing on Bandura's account of vicarious learning, argue that exposure to toxic behaviour as victim or bystander recalibrates perpetration, with the result that "toxicity can therefore be considered an acceptable and familiar gaming norm" (p. 428). Kordyaka et al. (2023) confirm the cycle empirically across 216 League of Legends and Dota 2 players: perpetrators, victims, and bystanders show highly significant positive correlations across all three roles (p. 397:3). The cycle does not require new perpetrators with hostile dispositions; it requires only that *exposure* recalibrate those who were initially neutral. Türkay et al. (2020) supply complementary player-side evidence from collegiate esports: accepting or avoiding, rationalising, and retreating were the dominant individual coping responses to witnessed toxicity, set within an overarching normalisation of toxic conduct as part of competitive culture (pp. 1, 7).

2.5. The Gendered Distribution of Costs

Any structural account must explain how harm is distributed. Fox and Tang (2017) demonstrate that female players face a qualitatively distinct form of toxicity: sexual harassment, but not general harassment, triggers rumination and produces withdrawal mediated by perceptions of organisational unresponsiveness, and "this apparent indifference leads to women quitting their games" (p. 1298). Female players' coping strategies, chiefly gender masking, are individually rational but collectively produce a spiral of silence as gender becomes unrecognisable (compare Fox & Tang, 2017). Wong and Ratan (2024) extend the picture: seeking help moderates the relationship between general harassment and anxiety, such that female gamers who seek help fare worse - a pattern the authors attribute to unfulfilled expectations of organisational response (pp. 681–682). Tang, Reer, and Quandt (2020) push the analysis to identity: gamer identification predicts sexual harassment perpetration, and the gamer identity is "associated exclusively to heterosexual masculinity" (p. 133), with the same logic extending to non-hetero identities. Kowert (2020) names the territorial logic: where games are framed as a "boy's toy," those who do not fit the mold are read as infiltrators, and harassment becomes a means of making them "leave their space," a logic Andéhn et al. (2024) trace across the wider manosphere. Sexual harassment functions as a mechanism for policing who belongs.

2.6. Governance Failure

Platform governance has not displaced this normative order, and players themselves recognise the gap: in a 2019 Anti-Defamation League survey, 62% of players said companies should do more to make online games safer and more inclusive (as cited in Kowert, 2020). Kou (2021) shows that permanent bans produce a "stereotype of 'the most toxic player'" rather than reform (p. 334:1); the punitive system "fails to account for the everyday toxic behavior by a 'good' account that often erupts out of intense negative emotions" (p. 334:16). The reporting infrastructure on which moderation depends is itself co-opted: Kou and Gui (2021) show that flagging "becomes a mechanism for players to cope with game loss, and a retaliatory act against the person responsible for the game loss." Reward systems suffer the same fate: a peer-endorsement instrument designed for prosocial conduct has been redirected into a recognition device for competitive performance, as a large gaming developer itself acknowledges (Poeller et al., 2023, p. 374:2). Muting the chat, the dominant coping mechanism by other players, places the burden of safety management on those least responsible and leaves perpetrators "unchallenged in their assumption that their behaviour is indeed normal" (Poeller et al., 2023, p. 374:21). Each governance instrument fails along the same axis: it processes individual conduct in a community whose hostility is produced collectively.

2.7. The Theoretical Gap and the Conversation This Thesis Joins

Research on hostile conduct and platform governance in competitive online multiplayer gaming establishes a robust structural account: definition, disinhibition, toxic meritocracy, normalisation through vicarious learning, gendered distribution, and governance failure. It says less about how community members construct the normative order that renders hostile conduct legitimate. Three gaps motivate this thesis. First, when researchers treat players' reasoning, they tend to read it as moral disengagement (e.g., Poeller et al., 2023, p. 374:5): a cover story for behaviour produced elsewhere, rather than the substantive content of a moral order. Second, the structural accounts work at scales (platform, culture, demographic) that explain why toxicity persists but say little about how players legitimise it match by match in their own words. Third, researchers have diagnosed governance failure without theorising it: where official moderation falls short, players develop their own customs of acceptable conduct, and these customs remain underspecified.

This thesis treats community members' rationalisations as the normative architecture itself: the moral order in which players produce, evaluate, and reproduce hostile conduct. Closing this gap requires a theoretical apparatus built for analysing customary moral orders rather than deviations from them. Thompson's customs in common, Polanyi's commodity logic, and Sayer's lay morality, synthesised by Bolton and Laaser and developed in the next chapter into a three-layer framework adapted to symbolic-stakes, pseudonymous, transient communities, supply that apparatus.

3. Theoretical Framework

This chapter develops moral economy as the thesis's analytical lens. The gaming literature accounts for two levels: how platform design produces hostility, and how individual players process it. It leaves the community level in between underdeveloped, where members collectively render hostile conduct legitimate. Moral economy theorises that level. Section 3.1 traces the framework from Thompson to Bolton and Laaser. Section 3.2 specifies its three layers and applies them to competitive gaming, with the loop in Figure 1. Section 3.3 sets out four ways competitive gaming differs from the framework's original setting, and the adaptations these differences require.

3.1. From Thompson to Bolton and Laaser

Moral economy is a framework for analysing the moral norms that communities produce, defend, and contest in the conduct of economic life. Thompson coined the term to make sense of eighteenth-century English food riots, dismissed by orthodox accounts as “rebellions of the belly” (Thompson, 1971, p. 77). Where those accounts read the riots as instinctive responses to hunger, Thompson re-reads them as actions in defence of “traditional rights or customs” (p. 78). For Thompson, *customs* are not paternalistic norms imposed from above; they are the community’s own standards, made and defended by its members. Outrage at violations of these moral assumptions, as much as material deprivation, was the usual occasion for direct action (p. 79). Within the community’s logic, the hostile responses that followed such violations were forms of legitimate enforcement, not aggression. *Customs in Common* (Thompson, 1991) generalises the move: rather than a stable inheritance, a custom is “a field of change and of contest, an arena in which opposing interests make conflicting claims” (p. 6); contemporary norm research treats violation and enforcement in similar dynamic terms (Van Kleef et al., 2019). If customs are made by the agency of those who hold them, which communities the framework can illuminate is an open question.

Bolton and Laaser (2013) propose a moral-economy synthesis that combines Thompson with two further authors, partly to address what they identify as the underdevelopment of individual reasoning in both Thompson’s and Polanyi’s moral economy (p. 515). From Karl Polanyi they take the structural side: his concept of *fictitious commodities* (land, labour, money), and his claim that the commodification of labour strips away its “deeply human and vulnerable character” (Bolton & Laaser, 2013, p. 512), and his “double movement” between market expansion and the protective social institutions that respond to it. From Andrew Sayer they take the individual side: his concept of *lay morality*, the everyday evaluative reasoning through which people judge “how people should treat one another in ways conducive to well-being” (Sayer, 2005, p. 951, as cited in Bolton & Laaser, 2013, p. 516). The combined claim is that political economy and ordinary moral reasoning have been studied apart, and the synthesis brings them back together. Lay morality “breathes life into a political economy framework” (Bolton & Laaser, 2013, p. 511) by populating structural accounts with reflective people who care, or fail to care, about each other.

The result is the three-layer framework (structural, community, individual) that this thesis adopts. Bolton and Laaser developed it for paid employment; Section 3.3 sets out the differences competitive gaming presents and the adaptations they require.

3.2. The Three Layers Applied to Competitive Gaming

The three layers are taken from Bolton and Laaser (2013, p. 509), who describe moral economy as forming an “analytical bridge” between individual agency, the institutionalised structures of community and work, and political economy. Polanyi sits at the structural level: he explains in this specific case how skilled play is being commodified by the platform. Thompson sits at the community level: he names the customs that organise legitimate conduct. Sayer sits at the individual level: he names the everyday moral reasoning through which conduct is justified or condemned. Figure 1 below illustrates how the three layers connect into a single reproductive loop.

The structural layer translates Polanyi’s commodity logic into platform architecture. Modern games organise players through ranked ladders: public hierarchies in which matchmaking sorts players by a skill rating updated each match. The system reduces skilled play to a single tradeable number that determines access, status, and recognition. Polanyi argues that pricing labour through exchange value strips away its “deeply human and vulnerable character” (Bolton & Laaser, 2013, p. 512); rating-based pricing produces a community in which a player’s perceived worth is inseparable from their ladder position. Gaming research has, as previously mentioned, named this dynamic toxic meritocracy (Paul, 2018, as cited in Kou, 2021, p. 334:5). The framework does not endorse the community’s view of merit. Matchmaking variance, queue partner quality, and network conditions all enter the rating, so “merit” is what the community reads into the number rather than something the number actually measures. What matters analytically is that the community treats it as fact. Kou and Gui (2021) document the consequence: players treat underperformance as flaggable in its own right, regardless of whether the player is also verbally toxic, and skilled but toxic teammates are tolerated where prosocial but unskilled ones are not. The platform’s commodity logic produces evaluative criteria that override conduct concerns.

The community layer translates Thompson’s customs into the unwritten match-conduct expectations every player is presumed to know: playing to win, accepting criticism from more skilled teammates, owing the rest of the lobby an honest effort. These specify what Thompson calls “the proper economic functions of several parties within the community” (Thompson, 1971, p. 79). When a custom is breached, the hostile response that follows is, within the community’s logic, enforcement rather than aggression. Sun et al. (2024, p. 428) identify vicarious learning as the mechanism through which “toxicity can therefore be considered an acceptable and familiar gaming norm” in the pseudonymous, transient groups Thompson did not envisage. Kordyaka et al. (2023, p. 397:3) find highly significant positive correlations across perpetrator, victim, and bystander roles, evidence that hostility circulates as community property. Platform governance,

premised on hostile conduct being deviance, fails when it confronts a community that treats the same conduct as enforcement. Kou and Gui (2021) record the absorption mechanism the framework predicts: the platform's reporting infrastructure, designed to flag misconduct, is redirected by the community into "a mechanism for players to cope with game loss, and a retaliatory act against the person responsible".

The individual layer translates Sayer's lay morality into players' real-time moral reasoning during and after matches (compare Bolton & Laaser, 2013, pp. 515–516): judging whether a teammate's mistake was forgivable, whether a callout was warranted, whether a report was justified or retaliatory. Poeller et al. (2023, p. 374:5) document this reasoning in the field, including the fatalism through which experienced players treat hostility as inherent to anonymous ranked play, and the moral-disengagement move that severs lay reasoning from accountability. Vicarious learning (Sun et al., 2024, p. 428) operates here as the transmission mechanism Thompson's slow sedimentation cannot supply in transient communities; it does so at a cost, since customs transmitted by exposure rather than by long shared history are more volatile. Each player remains a reflective person who can in principle accept or reject the customs (Bolton & Laaser, 2013, p. 515), yet the act of applying a custom to a specific incident reproduces the normative order the custom draws on. Customs are enacted into existence by lay moral reasoning, match by match.

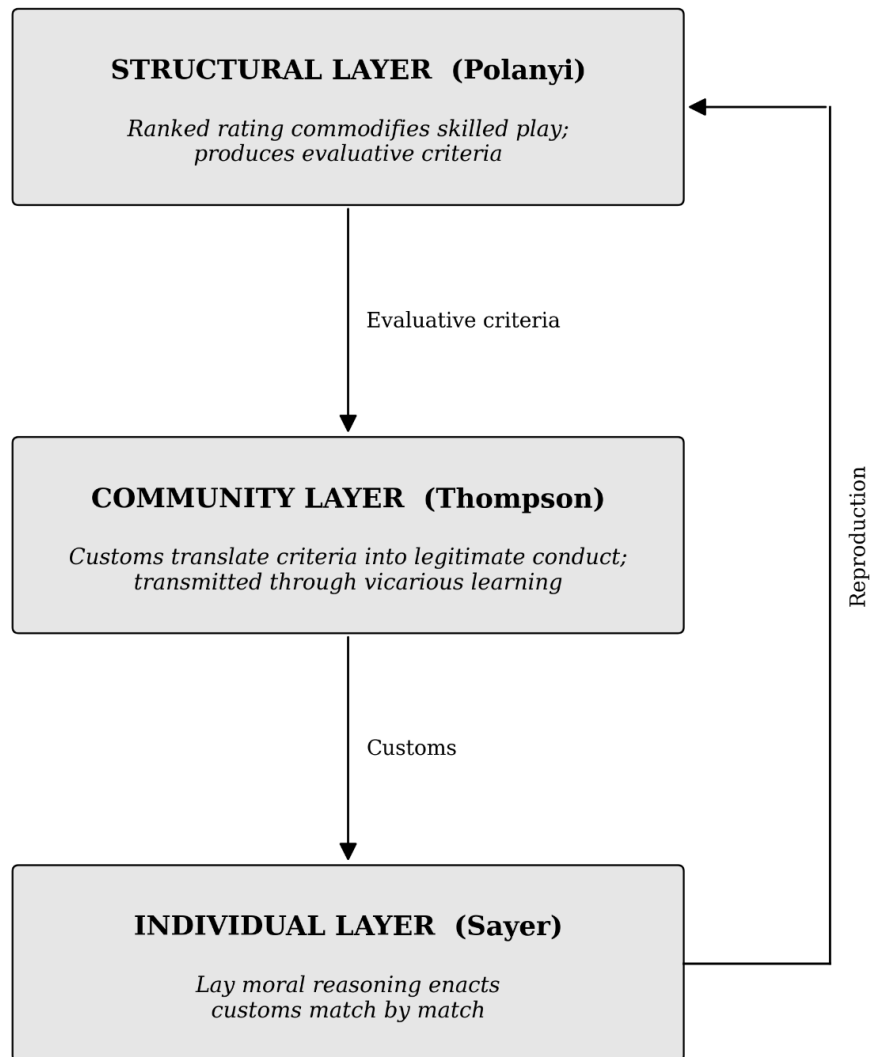


Figure 1. The three-level moral-economy loop, adapted from Bolton and Laaser (2013).

The figure reads top-to-bottom along the descending arrow and bottom-to-top along the return. The platform's commodity logic produces the evaluative criteria the community treats as the price of standing. Those criteria sediment into customs that the community defends through direct action. Lay moral reasoning enacts the customs match by match, and the individual decisions of who to defend, blame, or report feed back into the structural order. The loop closes when the customs' enactment becomes the platform's next input. Each layer is causally upstream of the layer below in the descending arrow, and each layer's reproduction depends on the layer above on the return. Bolton and Laaser developed this loop for paid employment; the next subsection sets out the four respects in which competitive gaming differs from that setting.

3.3. Adaptations

Moral economy as developed by Thompson, Bolton and Laaser, and Sayer was built for communities with material stakes and recurring face-to-face contact between known members. Competitive gaming differs from that setting in four respects, each requiring an explicit adaptation.

First, there is no employment relationship. In Polanyi's framework, labour's commodified status under market expansion makes workers, unions, and governments the constituency for the protective "double movement" (Bolton & Laaser, 2013, p. 512); workers are bound to that order through material need. In competitive gaming the stakes are symbolic (rank, reputation, standing), and players can in principle log off. But simultaneously disengagement carries a social cost: competitive multiplayer is a primary site of contemporary peer sociality, anchoring players to the community even when its conduct repels them. The thesis substitutes symbolic capital for material commodification and social need for material need as the binding force. What matters is that members are bound to an order they cannot freely exit, not which currency does the binding.

Second, there is no face-to-face community. Thompson's customs grew out of communities with stable membership and long shared histories (Thompson, 1991, p. 7-8); online matches are pseudonymous and transient, so the slow sedimentation Thompson describes cannot be the transmission mechanism. Transmission instead runs through the broader gaming culture and its vicarious-learning processes (Sun et al., 2024, p. 428). The cost, registered in 3.2, is that exposure-transmitted customs are more volatile than sedimented ones.

Third, pace. Thompson's framework presupposes generational depth in customary change; gaming norms shift on the timescale of patches. The framework's attention to long-run historical change must rebalance toward rapid customary turnover. Custom remains an arena of conflicting claims; only the cycle is faster.

Fourth, Sayer's account assumes "fellow-feeling" and "inter-dependence" between people in an ongoing relationship (Bolton & Laaser, 2013, p. 516). In competitive gaming, players often have no prior relationship and no anticipated future one. What functions as lay morality here is a thinner version of Sayer's, applied to strangers through community-transmitted norms. The form of lay reasoning persists; what is lost is the interpersonal warmth that grounds Sayer's account in the workplace, and the thesis recognises that loss.

4. Methodology

This study uses an interpretivist approach. It draws on twenty-two semi-structured interviews with active competitive online gamers and analyses the material through abductive thematic analysis, organised through a coding template. The focus is on how players themselves reason about hostile conduct: when it feels justified, when it crosses a line, and what the community

treats as acceptable. Behavioural and quantitative studies, which dominate the existing literature, can show how often hostility occurs but say less about why players see it as warranted (Bell et al., 2019, pp. 366, 376–377).

4.1. Research Approach

Interpretivism holds that social reality is shaped by the meanings people give to their actions (Bell et al., 2019, p. 31). If hostile conduct only makes moral sense within a community's shared norms, the best way to access those norms is through the players themselves. One risk is that participants present selective accounts of their own behaviour; cross-case comparison and independent coding help reduce this risk (Section 4.4).

The study links theory and data abductively, moving between interview material and theoretical ideas, and choosing the explanation that fits the data best rather than testing pre-set hypotheses (Bell et al., 2019, pp. 24–25). Existing structural accounts explain how hostile conduct spreads but say less about why players experience it as morally acceptable. The analysis tested several theoretical candidates against the emerging themes; moral economy was chosen when the data pointed to a gap between platform rules and what players themselves saw as legitimate.

Semi-structured interviewing followed from this approach. A pre-coded survey would impose categories on the very reasoning the study aims to capture; semi-structured interviews keep a consistent structure while allowing follow-up to follow the participant's framing (Bell et al., 2019, pp. 211, 435). Focus groups would have tended to produce shared, public-facing answers and would not work well given the disclosure risks discussed in Section 4.3. Thus interviews were conducted individually. Observation would record what players do but not why they consider it justified. A shared interview guide kept the conversations consistent across interviews (Section 4.3; Bell et al., 2019, pp. 438–440).

4.2. Sampling and Recruitment

The empirical setting is League of Legends, a team-based competitive online game with a large global player base. The toxicity literature has used this game often because its core features (a ranked ladder, public performance statistics, and pseudonymous matchmaking) are common across competitive online games more generally (Poeller et al., 2023). The aim is for the findings to apply to similar environments rather than to League of Legends alone.

The authors recruited twenty-two participants through purposive sampling, selecting on a clear criterion rather than aiming for statistical representativeness (Bell et al., 2019, pp. 391–393). The criterion was regular participation in ranked play, where the structural features the framework focuses on (rating, statistics, matchmaking) are most active and where the community norms the thesis examines are most clearly produced. Within this criterion, the authors prioritised long-tenured players, who tend to have a fuller grasp of those norms. Twenty-two interviews

allow comparison across cases while keeping the material at a manageable size for close reading (Bell et al., 2019, p. 397). See Appendix 3 for the full participant overview.

Recruitment combined direct outreach with snowball sampling (Bell et al., 2019, pp. 395–397). Snowball sampling tends to produce non-representative samples that reflect existing social networks rather than random variation. The final pool reflects this: it skews male, EU-Nordic, and English- or Swedish-speaking, with most participants having played for more than five years. This homogeneity is treated as a design limitation (Section 4.6).

4.3. Interview Guide and Data Collection

The guide asked participants to describe community norms in their own words before moving on to broader evaluation, on the basis that interview language should follow the participant's own usage where possible (Bell et al., 2019, p. 440). Questions moved from concrete experiences of play, to community norms around hostile conduct, to views on platform governance. This sequence follows the framework's three layers (individual, community, structural) and grounds broader judgements in events the participant has already described, which we judged would help reduce after-the-fact justification. Appendix 1 contains the full guide. Interviews took place via video call between 4 and 15 March 2026 and produced about thirty hours of recordings.

The disclosure issue in this study is unusual. Within the community's own norms, hostile conduct is often seen as legitimate; admitting to it in an interview can recast it as bad behaviour. Participants might therefore hold back about their own conduct, or about the emotional impact of hostility from others, since competitive culture often treats this as weakness. Each interview opened with the framing in Section A1.1, which presented both positive and negative experiences as valid, made clear that the researchers were familiar with the gaming community, and stressed that the goal was to understand rather than to judge.

Participants chose Swedish or English. Coding was done in the original language; only the quotes selected for the final write-up were translated. Translation involves rebuilding meaning across languages and cultures (Bell et al., 2019, pp. 450–451), and coding in the original language preserved the expressions the analysis depends on.

4.4. Analytical Procedure

The analysis followed an abductive logic, moving between data and theory rather than testing pre-set hypotheses (Bell et al., 2019, p. 24). With twenty-two transcripts and around thirty hours of recordings, an unstructured thematic approach would have been hard to manage. The authors used template analysis (Tabari, King, & Egan, 2020, pp. 199–200), which suits abductive work because the coding template is built early but is allowed to change as the analysis develops.

In the first round, both authors read the transcripts independently and coded openly, staying close to the participants' wording and noting anything evaluative: judgements about who deserved

blame, what counted as fair criticism, when a report was warranted. The structural accounts most common in the literature (online disinhibition, toxic meritocracy, vicarious learning) turned out to be too broad to organise the moral language that kept appearing. Pooling the codes across interviews produced a vocabulary of who deserves what, justified enforcement, and contested legitimacy that participants used when talking about reports, bans, and the behaviour behind them. Figure 2 summarises the result: eighteen focused codes consolidate into the six higher-order themes that organise the empirical findings chapter, all converging on the research objective.

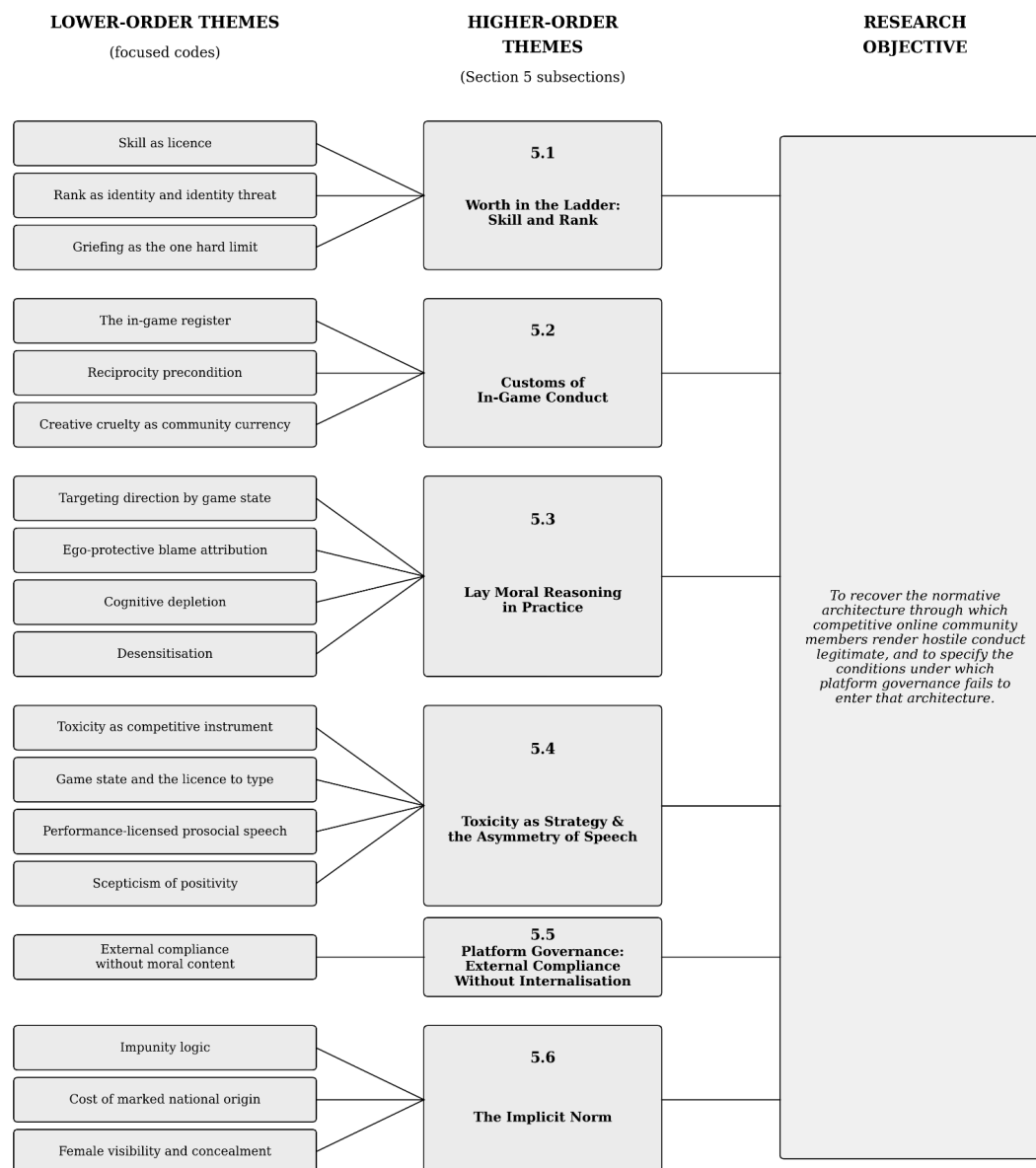


Figure 2. Summary of empirical findings: focused codes consolidated into higher-order themes (Section 5 subsections), mapped to the research objective.

Coding can break data into pieces and pull statements out of context (Bell et al., 2019, pp. 533–534). Open codes stayed close to the participants' wording; higher-order themes were done independently and then discussed thoroughly between authors.

The two researchers brought different positions to the study. One has substantial experience playing the game and is part of its community; the other is an adjacent outsider with experience of online gaming generally but less direct involvement in this specific community. Reflexivity matters in a particular way here (Bell et al., 2019, p. 156): the study's focus is the community's own moral language, which an insider can read more readily but can also take for granted. At the same time the adjacent outsider understands the community at large while needing to ask clarifying questions to understand the specifics. The authors used adversarial collaboration to manage this, an approach in which researchers who disagree resolve their differences through joint reasoning instead of compromise (Mellers et al., 2001). Each coded before they compared, treated disagreement as useful, and worked through ambiguous cases rather than splitting the difference. They also kept short debriefs after each interview as additional material.

4.5. Research Quality and Ethics

Research quality is assessed against Lincoln and Guba's four criteria (Bell et al., 2019, pp. 363–365). Credibility is supported by triangulation across twenty-two participants and by checking each higher-order theme against full transcript passages. Transferability is supported by detailed description of the setting, the sample (Appendix 3), and the structural features that affect when the findings should travel. Dependability is supported by an audit trail: discarded interpretations, and adversarial collaboration produced a written record of disagreements in various drafts. Confirmability is addressed through the insider-outsider split described in Section 4.4 and through the rule that any claim in the final write-up must be traceable by both authors to specific transcript passages.

The study followed the Stockholm School of Economics' ethical standards and the relevant GDPR rules (Bell et al., 2019, pp. 114, 118, 124–125). Participants gave informed consent and could withdraw at any point. Identifying details were removed and pseudonyms used throughout. Recordings and transcripts were stored on access-restricted devices. Where participants disclosed their own hostile conduct, the material is presented in analytical terms rather than direct quotes if attribution could risk identification.

4.6. Limitations

Four limitations qualify the design. First, network-based recruitment carries a risk of producing samples shaped by existing social networks rather than random variation (Bell et al., 2019, pp. 395–396); the toxicity literature itself is dominated by particular games and by

English-speaking, male-skewed samples, and the present sample falls within a similar demographic range. Second, recruiting through mutual acquaintances raises a contamination risk: participants who know each other socially may compare impressions after their interviews and adopt similar phrasings. Several recruitment channels (in-game requests to randomly encountered players, mutual acquaintances, social-media gaming communities) were used to avoid drawing all participants from a tight network. The risk is reduced rather than removed. Third, the analysis is interpretive and shaped by the authors' positions; the insider-outsider split and adversarial collaboration help limit confirmability (Bell et al., 2019, p. 365) without removing it. Fourth, the findings aim at transferability rather than statistical generalisability (Bell et al., 2019, p. 365), and are likely to apply where comparable environments share the structural features the framework focuses on.

5. Empirical Findings

This chapter presents interview material from twenty-two competitive online players, organised under six themes. Interpretation belongs in Chapter 6.

5.1. Worth in the Ladder: Skill and Rank

5.1.1. Skill as Licence

Many participants accepted that skill asymmetry licensed aggression. Connor put the position bluntly from the receiving side:

I'd much rather have a guy calling me a bunch of racial slurs but carrying me to victory. — Connor

Some, however, rejected the claim. Adrian said: "I don't think being good and carrying gives you any rights that might be toxic... you don't have any authority." Stella argued that a friendly player plays better for the team, and Tobias preferred a friendly but less-skilled teammate.

5.1.2. Rank as Identity and Identity Threat

For Felix, rank was identity:

You are very strongly tied to your rank. It almost becomes your identity in the game. And losing rank becomes almost a notch in the identity. — Felix

Several participants treated rank as a personal stake rather than a metric. Asked whether high-performing players were heard differently, Melody answered: "Yeah, definitely... [low performers' opinions] are not as valid and everyone can see it."

5.1.3. Griefing as the one hard limit

Most participants drew one hard limit on legitimate conduct: griefing, the deliberate sabotage of one's own team's match. Several described almost any verbal hostility as tolerable provided it stayed verbal. Blake stated the principle:

The chat is free – write whatever, just don't ruin the game. — Blake

Connor drew the same line from the receiving side: “not only have you lost the game for yourself, you're going to actively ruin the game for the other four people playing.” When asked what conduct they considered most serious, multiple respondents named griefing rather than verbal toxicity.

5.2. Customs of In-Game Conduct

5.2.1. The In-Game Register

Several participants described the in-game register as a separate space governed by its own expectations - a boxing ring, if you will. Ulric framed it explicitly:

When you sign up and enter a game, you should be prepared for people commenting on your gameplay. — Ulric

Melody offered the same framing from the position of someone subject to it: “when you are playing bad, you will get harsh treatment. And I think that's just part of the game.”

5.2.2. Reciprocity precondition

Connor articulated the boundary participants drew between banter and harassment:

Banter stops when it's no longer reciprocated... if you're trash talking someone and they're just like responding nicely, I think most honest, most normal people will just kind of stop. — Connor

Adrian named the precondition: “I would never say bad things to a full stranger that I just met. I need to actually like feel the water.”

5.2.3. Creative cruelty as community currency

Several participants described how creative cruelty was treated as a community competence. Hugo articulated the evaluation:

If you put a lot of effort into being hateful in a creative way... it's almost kind of admirable. — Hugo

Leon pointed to specific formulations recognised across the community: “Go on Amazon and buy rope, or hire garage space... it's classic.” The phrasing is a coded suicide-baiting formula. Felix described how platform restriction triggered linguistic adaptation: “They've changed chat

so you can't write cancer anymore. It used to be a classic that someone's mum would have cancer."

5.3. Lay Moral Reasoning in Practice

5.3.1. Targeting Direction by Game State

Several participants identified a pattern in the direction of hostility: when a team was losing, hostile chat tended inward toward teammates; when winning, outward toward opponents.

5.3.2. Ego-protective blame attribution

When describing losing matches, several participants articulated redirecting responsibility outward to preserve a sense of personal competence. Ulric named the mechanism:

Either when you have done something ineffectively, or if you have missed an opportunity, or just straight up lost, then I think that the toxicity is a way of coping and offloading the errors to other persons. — Ulric

The same protective move extended into platform tools. Connor described the in-game reporting function: "I don't think they're reporting you because they genuinely feel like you've been toxic... They're just doing it to ruin your game." Dominic put this more bluntly: "I love reporting people because I know they're going to get really mad... I'm reporting, not because I care about the community, but because I know that he's going to get mad from getting that ban." Roman articulated the underlying function:

I think it's more of an emotional coping function... you try to deceive yourself a little when you report people. — Roman

5.3.3. Cognitive depletion

Nico linked long sessions to a drop in self-restraint, using the analogy of a mental fuel tank being depleted:

When playing 16 hours in a row... the brain can't be bothered to think about what could be done differently, so it takes the shortcut and says it's everyone else's fault. — Nico

Several other participants made the same connection between cognitive depletion and lower thresholds for hostility.

5.3.4. Desensitisation

Several participants described how repeated exposure shifted what registered as significant. Tobias spelled out the threshold: "Someone can take, you know, some harassing for 30–40 minutes... that's fine. It's part of the game." Melody located the move at the level of culture rather than individual disposition:

I don't think that all the people that are typing racist or homophobic things are actually racists in real life. That's just the culture. — Melody

5.4. Toxicity as Strategy and the Asymmetry of Speech

5.4.1. Toxicity as Competitive Instrument

Ivan put hostility in competitive terms:

The thing is that you want to get inside the opponent's head, in a way – to maybe make them angry and make a few more mistakes. — Ivan

Multiple respondents described the move in similar terms.

5.4.2. Game State and the Licence to Type

Players described different thresholds for typing depending on whether their team was ahead or behind, with even games described as quieter games. Once a team had a clear lead, hostile chat opened up. Blake captured the asymmetry from the first-person side:

For me, I write more when I'm winning than when I'm losing. Writing when I'm losing feels like opening a door for the other team to walk through and step on me. — Blake

Several interviewees also described escalation against teammates whose play crossed into perceived sabotage.

5.4.3. Performance-licensed prosocial speech

Players who tried to defuse hostility or stand up for a targeted teammate described needing to be performing well themselves before speaking with any expectation of being heard. Julian set out the condition:

If I'm doing my job and we're winning, then I usually have the unspoken right to defend someone, because I'm doing something. There's nothing wrong with me, so to speak. But if I see that the person who made all the mistakes and is losing the match, then my tendency to protect someone would be lower. — Julian

Across the interviews, both high and low performers defaulted to muting as the practical solution.

5.4.4. The Scepticism of Positivity

Most participants described positive speech as something to be suspicious of. Dominic described the suspicion as a signalling problem:

If you're overly positive, then – in my head – it signals that the person is trying to annoy me. — Dominic

Melody, by contrast, had tried injecting positive chat into matches and described what happened in one instance: she wrote “well played” to a teammate she genuinely thought had played well, and the teammate, already tilted, read the message as sarcasm and responded with “kill yourself, trash” repeated about twenty times. Quentin offered a complementary diagnosis:

Because it feels a bit fake, firstly. But if you don’t have a lot of value in what you’re saying, it’s hard to trust what you’re saying. When you eventually do have something important to say, it’s just like ‘oh, you don’t trust them anymore’, because they said too much. — Quentin

5.5. Platform Governance: External Compliance Without Internalisation

Riot Games provides several conduct-shaping mechanisms: chat filters, a reporting tool, automated sanctions, and a peer-endorsement system. Asked whether his self-restraint rested on moral conviction or fear of sanction, Dominic answered:

It’s not like the moral that keeps you from writing things. It’s the afraid of getting banned in that way. — Dominic

5.6. The Implicit Norm

This section presents accounts from players whose experiences sit outside the community’s default participant.

5.6.1. Impunity Logic

Asked about hostile conduct directed at women, Ivan answered without prompting:

It’s because you can get away with it. You get no backlash whatsoever. It just doesn’t happen. — Ivan

5.6.2. Cost of marked national origin

Dominic plays under an in-game name with the suffix “PL” for Poland. He described the targeting that follows:

My in-game name has “PL” at the end, so usually they go say Nazi jokes... They’re not gonna get banned for it. Those sting the most for me. — Dominic

5.6.3. Female visibility and concealment

Melody concealed her gender during matches as a deliberate strategy: “I have an in-game name that doesn’t expose my gender... I would not want people to know that I’m a woman playing, so I don’t say it. They were talking about voice chat on League — I might not be a part of that. Even though it would be fun, I just feel like there’s just a lot of toxicity and less patience.”

Stella, who played as visibly female, recounted what that visibility attracted:

I have received extremely severe threats. I started a game with two friends and we got two randoms with us. For three games in a row we had the same two randoms. In the first game they realised I was a woman; in the second game one of them started describing how he would rape me; in the third game he described in detail how they would cut me open and have me hang myself with my own intestines. So it is genuinely extreme psychological straining. — Stella

6. Analysis

The customs the participants describe form a moral order, not a deviation from one. Read through Polanyi, Thompson, and Sayer, the data show a single mechanism operating across the framework's three layers. The platform's ranked rating turns skilled play into the currency that buys standing in the community, so the rating effectively prices voice. The community's customs then convert standing into a licence to enforce: high-rated players are entitled to discipline lower-rated ones who fall short of match-conduct expectations. Lay moral reasoning, finally, decides whom that licence falls on in any given match. The mechanism explains the general pattern of hostile conduct and what platform governance cannot do about it. Section 6.1 traces the customary loop through a normal match; 6.2 names what the order valorises, which is not what its loudest participants do; 6.3 returns to platform governance and shows why each instrument is absorbed; 6.4 reads gendered targeting as the loop's distributional residue.

6.1. Skill-as-Licence as Moral Enforcement

This theme operates across all three layers of the framework. Polanyi's commodity logic sits at the structural level, Thompson's customs at the community level, Sayer's lay morality at the individual level. The three layers form one customary loop: the rating supplies the criteria, the customs convert them into a licence to enforce, and lay reasoning enacts the licence match by match. The gaming literature treats hostile conduct as the problem to be explained: toxicity to be reduced, moral disengagement to be undone, governance failure to be repaired. The customs framing reads hostility as the order's normal output. Focused codes track the loop's stages: symbolic compression of worth onto the rating, performance-rationed standing, the in-game register, the reciprocity precondition, performance-licensed aggression, ego-protective blame, cognitive depletion, and desensitisation.

Every ranked match adjusts a number on the player's account; the community treats that number as the price of voice. Under Polanyi's commodity logic, skilled play takes a tradeable form: the rating. The participants supply the moral content the rating itself does not carry. Several in 5.1 described the rating as identity, the under-performer's opinion as invalid, and standing as a function of recent play. Felix put it plainly: "you are very strongly tied to your rank. It almost becomes your identity in the game." For the community, the number measures merit. Matchmaking variance, queue-partner quality, and network conditions all enter it; the community defends the resulting ascription as merit earned. Kou and Gui (2021) record the signature:

under-performance is itself flaggable, regardless of speech. Their data show what their reading does not name: the rating supplies the criterion the customs moralise.

Customs do what the rating cannot. Many participants in 5.2 described conduct expectations the community never states but expects every player to know. By entering ranked, the player accepts hostile commentary as part of the contest. This is custom in Thompson's specific sense (1991, p. 6): a field of contest defended through direct action in the in-game register, not the diffuse pressure of a modern norm. Two preconditions distinguish enforcement from arbitrary attack. Reciprocity: participants in 5.2.2 describe banter that ends when one side disengages, and players give strangers a closer reading before targeting. Skill-as-licence: rating earns commentary, commentary earns standing. Several participants in 5.1.1 framed the licence as moral entitlement, ranking the carrying player's right to slur above the friendly under-performer's right to be spared.

Lay moral reasoning enacts the customs match by match. Ego-protective blame is the customs' targeting apparatus, not a separate coping mechanism. After a loss, the player decides who deserved the report and who threw the game (community shorthand for losing a winnable match through bad play). That verdict supplies the moral content the report itself lacks. Without it, the player can press the button but cannot say what for. Several participants in 5.3 described reporting as retaliation against the teammate they had decided lost the match, not as flagging misconduct. Nico described long sessions as a failure mode: cognitive load collapses the apparatus, the player stops adjudicating, and the customary default applies. The rest of the team lost the game, not them. Poeller et al. (2023) read this reasoning as moral disengagement, a cover story for behaviour produced elsewhere. The framework reads it as enforcement: lay reasoning faithfully enacting the licence the customs supply.

6.2. Apex Restraint as the Order's Ideal

This theme operates at Thompson's community layer as a counter-current within the customs: the player who has the licence the rating gives and declines to use it. The thesis names this figure apex restraint, drawing on the focused code of the same name; participants describe the figure but do not theorise it. The community calls them the strong and silent type.

By the participants' own reports in 5.4, the most skilled player speaks little, with composure, in a register free of the insecurity that fuels most chat. The team listens. Participants describe the type with admiration but treat it as a fact about good players, not about the order. The framework reads it as a fact about the order: the customs reserve their highest standing for the figure who declines what the rating supplies.

The composure does the work. Hostile players seek a specific reward: the rush of having ruffled the target. The apex player denies it. Refusing to react starves the hostility of its payoff. Restraint reads as authority because the order recognises the counter-move.

The order's authority figure is the player who least uses the licence. The customs licence hostile speech by skill and reserve their highest standing for the player who declines. That contradiction gives the order its credibility. A community whose ideal was the loudest aggressor would defend its conduct as taste. A community whose ideal is the player who could attack and chooses not to defends its customs as morally serious. Accounts in 5.4 register the ambivalence: participants who themselves use the licence admire the player who does not. Toxic meritocracy (Paul, 2018, as cited in Kou, 2021, p. 334:5) names the status hierarchy the rating produces. The apex-restraint finding identifies what that account leaves out: the moral content that anchors the hierarchy, and the figure who carries it.

The scepticism of unprompted positivity and apex restraint are the same custom from two angles. The community distrusts unprompted praise because credibility comes from performance, and apex restraint clears that bar. Silence-with-skill is the only credible form of approval the order admits. Dissent against skill-as-licence has no customary architecture. The participants in 5.1.1 who reject the licence (Adrian on rights, Stella on team play, Tobias on the friendly teammate) register their objection but cannot convert it into communal force. The apex-restraint figure operates with the architecture they lack. The most credible figure in the order does not undermine it.

6.3. The Order Absorbs Governance

The literature reads each platform instrument as an enforcement failure. Kou (2021) shows that permanent bans produce a stereotype of the toxic player rather than reform; Kou and Gui (2021) show the report tool becomes retaliation against teammates blamed for losses; Poeller et al. (2023) show peer-endorsement has been redirected toward competitive performance. The framework reads the same evidence differently. Gaming developer's structural instruments arrive on terrain Thompson's customs have populated. Each enters with its meaning fixed by the customs it was built to displace. The customs absorb the instrument.

The reporting tool, designed to flag misconduct, becomes a weapon teammates use against the player they decided lost the match (5.4, 5.5); the customs have supplied the moral content authorising retaliation; the tool inherits it. The chat filter generates puzzles, and players valorise the workaround as competence (5.2); the patch cycle turns Thompson's slow rhythm over fast. The peer-endorsement system asks players to praise teammates in a community that treats unprompted positivity as untrustworthy (5.4.4); the only approval the customs admit is the absence of attack, which the order names apex restraint (6.2). Muting, the dominant coping move, registers in the loop as the absence of a problem (5.5). The targeted player exits, the perpetrator continues, the order persists.

The pattern is self-rinsing in the wrong direction. A banned perpetrator re-registers and resumes; a targeted player who leaves often leaves for good. The exclusion falls on those it should protect, here women and players whose names mark national origin (5.6); the loudest flow back in.

The literature reads four distinct failures. The framework reads one mechanism: each processes individual conduct in a community whose hostility is produced collectively. The tools secure external compliance; the customs supply the moral content enforcement was supposed to install.

6.4. The Same Logic, Differently Distributed

Gendered-toxicity research reads female targeting as a separate mechanism: Fox and Tang (2017) as sexual harassment layered onto hostility, Tang, Reer, and Quandt (2020) as a gamer identity bound to heterosexual masculinity, Kowert (2020) as territorial logic framing women as infiltrators. The framework reads it as one loop's distributional residue: 6.1's enforcement logic, redirected wherever the customs decide a target produces no backlash.

Ivan supplies the principle: *"It's because you can get away with it. You get no backlash whatsoever"* (5.6.1). The custom of justified enforcement needs a target the community will not defend: under-performance codes enforcement as performance-driven (6.1), visible femininity (5.6.3) or origin-marked names (5.6.2) as gender- or origin-driven. Misogynistic and xenophobic affect appear but do not exhaust the explanation; targeting follows the absence of community-internal accountability, which selects where the loop lands.

The data establish a severity asymmetry. Stella received sustained rape and lethal threats at a register no other target category attracted (5.6.3); Melody concealed her gender (5.6.3); Dominic took targeting attached to his "PL" suffix as what stings most (5.6.2). A nail that stands out will be hammered down: any deviation from the generic pseudonym supplies identification, and the creative-cruelty custom (5.2.3) takes over. Reaction-seeking promises a target made to flinch. Slurflation strips the edge from the old classics ("hope your mom gets cancer", "get a rope from Amazon"), and the customs reward whatever still lands. A Polish suffix opens Nazi jokes: sensitive enough to find the chink in the armour, new enough to demonstrate competence; visible femininity opens gendered violence on the same logic. Players pile into the slurfest the marker invites. Evading chat filters is a cat-and-mouse game within the game (6.3). Different markers, same logic, asymmetric cost: women and origin-marked names absorb what others do not. Impunity selects where the loop lands; creative cruelty selects what gets said. The gendered distribution follows from both.

7. Discussion and Conclusion

7.1. Answering the Research Question

This thesis asked how members of competitive online communities rationalize hostile conduct as normatively legitimate, and what that pattern reveals about platform governance failure. Players rationalize hostility through customs, community-internal rules enforced by participants: the rating prices voice, customs convert skill into a licence to enforce, and lay reasoning targets each match's enforcement. Neither the moral-disengagement reading (Poeller et al., 2023) nor the

under-enforcement reading (Kou, 2021) accounts for the architecture the participants describe. The insider-outsider split (4.4) shaped what the analysis read: the insider read the licence, the outsider read the architecture.

7.1.1. Sustaining the Moral Order in Competitive Online Play

Four mechanisms sustain the customs. Apex restraint: players reserve the highest standing for the high-skill participant who has the licence to flame and declines it. The absence of a public form for dissent: rejection registers as personal preference; no counter-figure emerges. Absorption: reports become retaliation, the filter prompts workarounds, endorsement recognises carry performance, mute hides hostility from the muter. Throughput: banned perpetrators return; targeted players leave for good.

7.2. Contribution to the Academic Debate

Bolton and Laaser developed moral economy for paid employment, with material need and stable membership doing the binding; the thesis extends it into pseudonymous, symbolic-stakes communities where vicarious learning replaces slow sedimentation and customs turn over with the patch cycle (3.3). The platform-governance literature reads each tool as an enforcement failure; the thesis reads each tool as entering a populated moral order that already supplies the moral content the tool was supposed to install (2.6, 6.3). Poeller et al. (2023) read player reasoning as moral disengagement; the thesis reads it as the order's substantive content, with evaluative consistency about who deserves to enforce (2.7). Paul (cited in Kou, 2021) names toxic meritocracy without naming the figure that anchors it; the thesis names apex restraint, and locates the order's credibility in the player who holds the licence and declines it (6.2). The gendered-toxicity literature reads female targeting as a separate misogyny mechanism; the thesis reads it as the loop's distributional residue, one impunity logic redirected wherever the customs decide a target will not produce backlash (6.4).

7.3. Implications for Practitioners and Management

7.3.1. For platforms and game studios

Each platform instrument enters a populated moral order that has already supplied the meaning the instrument was supposed to install (6.3). Studios can build apex restraint into a designed affordance: credit visible restraint by high-skill players, surface low-chat high-skill behaviour the way kill-death ratios are surfaced, rebuild peer-endorsement to recognise protective acts, and centre top-rated streamers who practise restraint in community programmes (6.2, 7.5).

7.3.2. For community managers and moderation system designers

Reporters use the report tool to retaliate against the teammate they decided lost the match (6.1, 6.3), and Kou and Gui (2021) document the pattern at scale. Moderation designers can weight reports along two axes. By reporter rank: a report filed by a player still in the loss-state carries less evidential weight than one filed cold. By reporter-target asymmetry: teammate-on-teammate reports filed within a losing match count differently from reports across matches.

7.3.3. For players outside the default

Report-and-mute individualises a structural problem; muting registers as absence rather than complaint, the perpetrator continues, and concealment becomes the practical response for marked players (6.3, 6.4). Platforms can give marked players pseudonymity controls that hide voice and identity-marker without disclosing why, route gender-targeting through a separate response pathway with transparency reporting, and recruit moderation staff from outside the demographic core to detect targeting the community will not defend.

7.3.4. Ethical implications

Treating the customs as a moral order risks naturalising them, and inaction transfers cost: banned perpetrators return, targeted players leave for good (6.3). Practitioners should treat the customs as the dominant order rather than the only one (Thompson, 1991, p. 6; 3.1, 6.2), and weigh interventions against the cost the order's victims pay; SDG 5 frames the gendered distribution, and SDG 16 frames the institutional gap platforms occupy as governors of speech and belonging for billions.

7.4. Limitations

Section 4.6 sets out the methodological limitations. Three further qualifications stand: the case is a single game, the lobby is the vantage rather than platform-wide production, and patches turn the customs over from cycle to cycle.

7.5. Future Research

Three studies follow. Test the framework where identity-binding regimes remove pseudonymity. Test whether apex restraint can be built into a designed affordance, with a major studio. Recruit a female-majority sample to clarify how women participate in producing the customs rather than only receiving them.

7.6. Conclusion

The rating prices voice, the customs convert skill into a licence to enforce, lay reasoning targets that licence match by match, and the order reserves its highest standing for the player who declines it. Platform tools enter this populated moral order and inherit the meaning it supplies.

Poeller et al. (2023) read player reasoning as moral disengagement; the data read it as enforcement. The under-enforcement literature reads each tool as failing; the data read each as absorbed. The gendered-toxicity literature reads female targeting as a separate misogyny mechanism; the data read it as one impunity logic redirected. For theory, moral economy travels into pseudonymous, symbolic-stakes communities, and toxic meritocracy gains apex restraint as the figure that anchors it (7.2). For practice, instruments must reach the customs themselves, and the cost of inaction falls on the players the order has decided will absorb it (7.3). The order will not be governed out of existence, only reshaped from within.

References

- Achterbosch, L., Vamplew, P., & March, E. (2024). Assessing the impact of grieving in MMORPGs using self-determination theory. *Computers in Human Behavior*, 161, Article 108388. <https://doi.org/10.1016/j.chb.2024.108388>
- Andéhn, M., Hietanen, J., & Wickström, A. (2024). Becoming red-pilled: Affective production in online countercultural collectives. *New Media & Society*. Advance online publication. <https://doi.org/10.1177/14614448241305420>
- Anti-Defamation League. (2025). *Playing with hate: How online gamers with diverse identity usernames are treated*. ADL Center for Technology and Society. <https://www.adl.org/resources/press-release/hate-and-harassment-present-almost-half-online-multiplayer-gaming-sessions>
- Bell, E., Bryman, A., & Harley, B. (2019). *Business research methods* (5th ed.). Oxford University Press.
- Bolton, S. C., & Laaser, K. (2013). Work, employment and society through the lens of moral economy. *Work, Employment and Society*, 27(3), 508–525. <https://doi.org/10.1177/0950017013479828>
- Fox, J., & Tang, W. Y. (2017). Women's experiences with general and sexual harassment in online video games: Rumination, organizational responsiveness, withdrawal, and coping strategies. *New Media & Society*, 19(8), 1290–1307. <https://doi.org/10.1177/1461444816635778>
- Frommel, J., Johnson, D., & Mandryk, R. L. (2023). How perceived toxicity of gaming communities is associated with social capital, satisfaction of relatedness, and loneliness. *Computers in Human Behavior Reports*, 10, Article 100302. <https://doi.org/10.1016/j.chbr.2023.100302>
- Icon Era. (2025). *In-game purchase spending habits statistics (2025)*. Icon Era. <https://icon-era.com/blog/in-game-purchase-spending-habits-statistics-2025.343/>

Icon Era. (2026). *Gaming industry statistics and revenue data 2026*. Icon Era.
<https://icon-era.com/statistics/gaming-industry-revenue-statistics/>

Kordyaka, B., Jahn, K., & Niehaves, B. (2020). Towards a unified theory of toxic behavior in video games. *Internet Research*, 30(4), 1081–1102. <https://doi.org/10.1108/INTR-08-2019-0343>

Kordyaka, B., Laato, S., Jahn, K., Hamari, J., & Niehaves, B. (2023). The cycle of toxicity: Exploring relationships between personality and player roles in toxic behavior in multiplayer online battle arena games. *Proceedings of the ACM on Human-Computer Interaction*, 7(CHI PLAY), Article 397. <https://doi.org/10.1145/3611043>

Kordyaka, B., Karaosmanoglu, S., & Laato, S. (2025). Defining toxicity in multiplayer online games: A systematic literature review. *Computers in Human Behavior Reports*, 19, Article 100698. <https://doi.org/10.1016/j.chbr.2025.100698>

Kou, Y. (2020). Toxic behaviors in team-based competitive gaming: The case of League of Legends. In *Proceedings of the Annual Symposium on Computer-Human Interaction in Play (CHI PLAY '20)* (pp. 81–92). ACM. <https://doi.org/10.1145/3410404.3414243>

Kou, Y. (2021). Punishment and its discontents: An analysis of permanent ban in an online game community. *Proceedings of the ACM on Human-Computer Interaction*, 5(CSCW2), Article 334. <https://doi.org/10.1145/3476075>

Kou, Y., & Gui, X. (2021). Flag and flaggability in automated moderation: The case of reporting toxic behavior in an online game community. In *Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems (CHI '21)*. ACM.
<https://doi.org/10.1145/3411764.3445279>

Kowert, R. (2020). Dark participation in games. *Frontiers in Psychology*, 11, Article 598947. <https://doi.org/10.3389/fpsyg.2020.598947>

Lapidot-Lefler, N., & Barak, A. (2012). Effects of anonymity, invisibility, and lack of eye-contact on toxic online disinhibition. *Computers in Human Behavior*, 28(2), 434–443. <https://doi.org/10.1016/j.chb.2011.10.014>

Liu, Y., & Agur, C. (2023). "After all, they don't know me": Exploring the psychological mechanisms of toxic behavior in online games. *Games and Culture*, 18(5), 598–621. <https://doi.org/10.1177/15554120221115397>

Mellers, B., Hertwig, R., & Kahneman, D. (2001). Do frequency representations eliminate conjunction effects? An exercise in adversarial collaboration. *Psychological Science*, 12(4), 269–275. <https://doi.org/10.1111/1467-9280.00350>

- Newzoo. (2025). *Global games market report 2025: Free version*. Newzoo.
<https://newzoo.com/resources/trend-reports/newzoo-global-games-market-report-2025-free-version>
- Paul, C. A. (2018). *The toxic meritocracy of video games: Why gaming culture is the worst*. University of Minnesota Press. (As cited in Kou, 2021.)
- Poeller, S., Dechant, M. J., Klarkowski, M., & Mandryk, R. L. (2023). Suspecting sarcasm: How League of Legends players dismiss positive communication in toxic environments. *Proceedings of the ACM on Human-Computer Interaction*, 7(CHI PLAY), Article 374.
<https://doi.org/10.1145/3611020>
- Sayer, A. (2005). Class, moral worth and recognition. *Sociology*, 39(5), 947–963.
<https://doi.org/10.1177/0038038505058376>
- Suler, J. (2004). The online disinhibition effect. *CyberPsychology & Behavior*, 7(3), 321–326.
<https://doi.org/10.1089/1094931041291295>
- Sun, X., Yu, V., & Chen, V. H. H. (2024). Toxic behavior in multiplayer online games: The role of witnessed verbal aggression, game engagement intensity, and social self-efficacy. *Chinese Journal of Communication*, 18(4), 426–444. <https://doi.org/10.1080/17544750.2024.2425662>
- Tabari, S., King, N., & Egan, D. (2020). Potential application of template analysis in qualitative hospitality management research. *Hospitality & Society*, 10(2), 197–216.
https://doi.org/10.1386/hosp_00020_1
- Tang, W. Y., Reer, F., & Quandt, T. (2020). Investigating sexual harassment in online video games: How personality and context factors are related to toxic sexual behaviors against fellow players. *Aggressive Behavior*, 46(1), 127–135. <https://doi.org/10.1002/ab.21873>
- Thompson, E. P. (1971). The moral economy of the English crowd in the eighteenth century. *Past and Present*, 50(1), 76–136. <https://doi.org/10.1093/past/50.1.76>
- Thompson, E. P. (1991). *Customs in common*. Merlin Press.
- Türkay, S., Formosa, J., Adinolf, S., Cuthbert, R., & Altizer, R. (2020). See no evil, hear no evil, speak no evil: How collegiate players define, experience, and cope with toxicity. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems (CHI '20)* (pp. 1–13). ACM. <https://doi.org/10.1145/3313831.3376191>
- Twinfinite. (2025, April 14). *Most played games in 2025 — ranked by average monthly players*. Twinfinite. <https://twinfinite.net/features/most-played-games/>

Van Kleef, G. A., Gelfand, M. J., & Jetten, J. (2019). The dynamic nature of social norms: New perspectives on norm development, impact, violation, and enforcement. *Journal of Experimental Social Psychology*, 84, Article 103814. <https://doi.org/10.1016/j.jesp.2019.05.002>

Wong, S., & Ratan, R. (2024). How does exposure to general and sexual harassment relate to female gamers' coping strategies and mental health? *Games and Culture*, 19(5), 670–689. <https://doi.org/10.1177/15554120231177600>

Zsila, Á., Shabahang, R., Aruguete, M. S., & Orosz, G. (2022). Toxic behaviors in online multiplayer games: Prevalence, perception, risk factors of victimization, and psychological consequences. *Aggressive Behavior*, 48(3), 356–364. <https://doi.org/10.1002/ab.22023>

Appendix 1. Interview Guide

The semi-structured guide reproduced below is the final iteration used across the twenty-two interviews. It is organised in five sections matching the inside-out movement described in Section 4.3: from the opening framing through individual context and community customs to platform governance, closing with reflection. Probes followed participants' own vocabulary and were not pre-scripted; not every question was put to every participant, and order varied with the flow of conversation.

A1.1. Opening and Framing

The verbal framing below was delivered conversationally at the opening of each interview, in the language of the interview (Swedish or English), and not read verbatim from a script. Its purpose was to establish conditions of psychological safety and to reduce social desirability effects across both directions of disclosure identified in Section 4.3.

Thank you for agreeing to participate in this interview. Before we begin, there are a few practical points we want to confirm with you.

On one hand, some of what we discuss may touch on moments where you yourself acted in ways that felt toxic, where you said things that were not politically correct, where you were overly emotional, overreacted, or simply were not at your best. On the other hand, it might go the other direction: you may not feel comfortable admitting that certain things offended you, affected you emotionally, or got under your skin more than you would like to let on. Both are completely understandable, and both are exactly the kinds of things we want to be able to talk about openly.

We also want you to know that we are gamers ourselves. We have seen how this world works. We know what goes on. So nothing you say is going to surprise or shock us.

Most importantly, we do not moralise. We are not here to judge whether something was right or wrong, good or bad. What we genuinely want to understand is the why – why certain things

happen, why people think and act the way they do in these situations. That is all. No judgments, no conclusions about who you are as a person.

So please feel free to speak openly. This is a safe space. Everything you share is anonymous and will be handled in full accordance with GDPR regulations and the ethical guidelines of the Stockholm School of Economics. You are free to skip any question you are not comfortable answering, and you may end the interview at any time without consequence.

Finally, do you consent to the interview being recorded? The recording will be used solely for transcription purposes and will not be shared outside the research team.

A1.2. Background and Entry into Competitive Play (Individual Layer)

1. When did you start playing League of Legends, and how often do you play today?
2. What game modes do you usually play: ranked, normal games, ARAM?
3. Do you usually play alone or together with friends?
4. What role or position do you usually play?
5. What originally motivated you to start playing?
6. How important is the game in your daily life or routine today?
7. Do you engage with the game outside of playing: streams, esports, Reddit, forums?
8. How would you describe your overall experience with the League of Legends community so far?

A1.3. Match Conduct and Community Customs (Customs Layer)

9. When you hear the term “toxic behaviour,” what does it mean to you?
10. What typically triggers frustration or conflict during a game?
11. In your experience, where do players draw the line between banter and behaviour that becomes genuinely toxic?
12. Are there situations where flaming or insulting someone feels justified to you or others?
13. Do you think there is more toxicity within your own team or toward the enemy team?
14. How do you react to an overly friendly player?
15. When something goes wrong and players start blaming each other, how do they usually decide who is at fault? What feels fair or unfair in those situations?
16. When someone behaves badly, how do teammates usually react: do they correct it, ignore it, report it, or respond in other ways?
17. Do players feel they owe something to their teammates during a match: patience, support, or effort?
18. Do expectations about behaviour change depending on the situation, for example ranked versus normal games, or strangers versus friends?

19. What does “normal behaviour” look like in League? What would you say is the group norm?
20. What behaviour makes someone a “good player” in the community?
21. What behaviour earns social approval from others?
22. What behaviour leads to ridicule or exclusion?
23. Is being skilled more important than being respectful, or vice versa?
24. Would you prefer a very skilled but toxic teammate or a friendly but weaker player?
25. If a high-performing player behaves aggressively, what effect does that have on the team?
26. Are players ever criticised for being “too sensitive” or “too aggressive”? In which situations?
27. What are the building blocks of a winning team in your view?
28. Which groups of players do you think are most at risk of being targeted by harassment?
29. Do you think it is more acceptable to be toxic in a fun or creative way?
30. Do players change their behaviour based on how others act in a match? Why or why not?

A1.4. Platform Governance: Reporting, Peer-Endorsement System, Sanctions (Structural Layer)

31. What does the peer-endorsement system represent to you?
32. What kind of player does the system seem to reward?
33. When someone receives endorsement after a game, what do you think that actually means?
34. Do you see endorsement as meaningful recognition or more as a technical feature of the game?
35. Does receiving endorsement influence how you see yourself as a player?
36. Do you think other players take endorsement seriously? Why or why not?
37. Does endorsement reflect genuine behaviour, or can it be gamed?
38. Does endorsement align with what the community actually values?
39. Do you think the peer-endorsement system is fair? Why or why not?
40. Do general players think Riot’s systems — reporting, chat restrictions, peer endorsement — are fair or unfair? Why?
41. Does public recognition change how people act?
42. Is prosocial behaviour visible enough to actually influence others?
43. Can recognition-based systems genuinely change norms, or do they only signal values?
44. Do you think punishment or reward is more effective in shaping behaviour?
45. If the peer-endorsement system disappeared tomorrow, would anything change?
46. Who is responsible for maintaining a healthy community atmosphere: players or Riot Games?
47. How do you think the peer-endorsement system could be improved?

A1.5. Reflection and Close

1. Is there anything I should have asked but did not?
2. What do you think this research project should pay most attention to?
3. Is there a moment from your time playing this game that has stayed with you?

Appendix 2. Interview Specifications

The full participant table with pseudonym, interview date and duration in minutes.

Table A3. 1: Interview specifications. All twenty-two participants. Total recorded material: approximately 30 hours.

No.	Pseudonym	Date	Duration (min)
1	Adrian	03-04	70.4
2	Blake	03-04	56.4
3	Connor	03-05	67.1
4	Dominic	03-05	76.9
5	Elias	03-05	99.1
6	Felix	03-06	99.3
7	Gavin	03-06	81.4
8	Hugo	03-07	78.3
9	Ivan	03-07	64.4
10	Julian	03-07	95.1
11	Kai	03-08	92.7
12	Leon	03-08	43.6
13	Melody	03-09	76.9
14	Nico	03-09	151.8
15	Orion	03-10	131.2
16	Pascal	03-10	93.1
17	Quentin	03-08	55.3
18	Roman	03-10	98.2
19	Stella	03-12	78.5
20	Tobias	03-13	43.1
21	Ulric	03-14	104.8
22	Viktor	03-15	81.4

Summary statistics. Average interview duration: 83.5 minutes (1:24). Minimum duration: 43.1 minutes (0:43, Tobias). Maximum duration: 151.8 minutes (2:32, Nico). Total recorded material: approximately 30 hours and 30 minutes across the twenty-two interviews. All interviews were conducted via video call between 4 and 15 March 2026.

Appendix 3. Use of AI Tools

AI tools were used for specific stages of this study.

The authors built a custom retrieval database from the corrected interview transcripts. The system accepts a paraphrased description of a remembered passage and returns the matching transcript location, with timestamp to the second. This compressed navigation across the roughly thirty hours of recorded material from manual searching to under thirty seconds per query. Every passage surfaced through the database was then verified by the authors against the original audio and the corrected transcript before any use in Chapter 5.

An automated speech-to-text tool produced the initial transcripts, which both authors corrected manually against the recordings.

A general-purpose language model was used for sentence-level polishing of the authors' own prose, without adding to or altering the content of the research. No AI tool was used to generate codes, interpret data, or produce analytical claims.