

Reputational Escalation in Post-Publication Peer Review: A Critical Discourse Analysis of the PubPeer Commentary on *'An Expert Criticism on Post-Publication Peer Review Platforms: The Case of PubPeer'*

Edward Mason

Independent Scholar, Publication Ethics and Research Integrity

Abstract

The PubPeer discussion of *'An expert criticism on post-publication peer review platforms: the case of PubPeer'* presents an unusually reflexive dispute in which a platform criticised by an article also hosted, moderated, archived, and partly answered the public response to that article. This commentary analyses a frozen record of 64 numbered comments posted under 32 visible display names, including pseudonymous contributors, named contributors, author responses, and two embedded moderator interventions. The analysis combines documentary verification with critical and computer-mediated discourse analysis, while preserving a strict distinction between reproducible article-level defects and unproved person-level conclusions. Several bibliographic mismatches and repeated passages can be checked directly against the published record, whereas claims concerning comprehensive artificial-intelligence generation, motive, bias, competence, collective responsibility, or personal integrity are not established by those observations. Across the thread, cautious source checking develops into categorical attribution, rhetorical questions carry presuppositions of failure, sarcasm and idiom recruit audience alignment, numerical lists convert publication events into personal tallies, and links to prior publications expand an article-specific discussion into a reputational dossier. The platform's moderator voice also adopts dismissive and self-credentialling language that is difficult to reconcile with PubPeer's published requirement that comments remain polite, neutral, and factual. The thread therefore cannot be treated as an adjudicative record merely because it contains some correct observations. Journals and institutions must independently verify provenance, relevance, expertise, conflicts, materiality, responsibility, and current claim status before allowing platform commentary to produce consequential judgements.

Keywords: PubPeer; post-publication commentary; critical discourse analysis; reputational escalation; anonymous commentary; moderation; epistemic authority; institutional non-delegation

Authorial position and evidentiary boundary

This commentary does not seek to validate, improve, promote, or legitimise PubPeer, post-publication comment websites, scientific watchdogs, or the individuals and organisations associated with them. It subjects their claimed competence, authority, methods, evidence, ethics, accountability, and institutional influence to sustained critical scrutiny. The analysis nevertheless observes an equally important boundary: it does not invent evidence, assume coordination, attribute hidden motives, or declare any participant

disqualified before the public record supports that conclusion. Correct article-level observations remain correct when made through a procedurally weak forum, while procedural weakness does not make an underlying observation false. The issue examined here is whether the discourse stayed within the evidentiary reach of those observations and whether the resulting thread possesses any justified claim to adjudicative authority.

1. Introduction: when criticism becomes its own case study

The article at the centre of this commentary criticised PubPeer and related post-publication platforms, after which PubPeer became the principal public venue in which the article itself was examined [1,2]. The resulting thread is analytically important because it collapses roles that would ordinarily be kept distinct. PubPeer was simultaneously the object of the published criticism, the technological host of the reply, the moderator of permissible language, the durable archive of allegations, and, through two moderator insertions, an institutional participant defending its own model. That configuration does not prove partiality, although it raises a demanding question about whether a platform can claim neutral authority while controlling the discursive conditions under which criticism of itself is received.

A rapid reading of the thread may produce a simple story in which commenters discovered defective references, authors responded, and the community corrected the record. Such a summary is inadequate because it treats every contribution as if it carried the same evidentiary status and remained directed at the same object. The thread begins with source-level observations, including claims that particular references or digital object identifiers could not be verified, but it subsequently moves through speculation about authorship technology, questions about whether the authors read their own text, numerical claims about retractions, assertions of bias and conflict, and increasingly personal characterisations. Critical discourse analysis is required precisely because the reputational effect lies in that movement from document to person rather than in any isolated sentence.

The distinction matters beyond one article because post-publication commentary can be searchable, persistent, indexed by general search engines, circulated through news and social media, and presented to editors or institutions as a ready-made evidentiary record. PubPeer's own frequently asked questions state that comments should rely on facts, logic, and publicly verifiable information and should remain polite, neutral, and factual, while the same guidance expressly states that PubPeer does not scientifically review comments and gives no warranty concerning their truth [3]. A thread may therefore possess high public visibility while lacking the scientific verification, identified expertise, conflict assessment, notice, reasoned determination, and appeal mechanism ordinarily expected before institutions impose consequential judgements.

This commentary asks four connected questions: which concerns in the thread can be independently reproduced; how did the wording change as the discussion proceeded; how did

repetition, quotation, links, lists, and moderator voice distribute credibility and blame; and what should journals or institutions infer from such a record? The purpose is not to rescue the focal article from verifiable criticism, because doing so would weaken rather than strengthen the analysis. The purpose is to demonstrate that accurate source checking and unjustified reputational escalation can coexist, and that the former cannot silently legitimise the latter.

2. Corpus, method, and analytic safeguards

The corpus is a frozen textual snapshot of the PubPeer thread associated with the focal article, comprising 64 numbered comments posted under 32 visible display names [2]. Sixty-three comments were dated November 2025, and one later update was dated January 2026. The displayed names included biological pseudonyms, several named contributors, and seven responses posted under the name Mohammad Abdollahi; two moderator interventions were embedded within comments rather than assigned separate comment numbers. The snapshot rendered some attached materials only as the word 'file', so this study analyses the public text and links but does not claim to have reconstructed every image or attachment originally visible to participants.

The method combines Fairclough's account of discourse as a relation among textual form, discursive practice, and social power with the broader critical-discourse framework described by Wodak and Meyer [4-6]. Herring's computer-mediated discourse analysis informs attention to platform structure, sequencing, uptake, pseudonymous identity, reply patterns, and the technical persistence of the exchange [7]. Appraisal theory supports analysis of judgement, affect, graduation, and evidentiary commitment, while Culpeper's work on impoliteness assists in distinguishing technical disagreement from language that attacks competence, social standing, or face [8,9]. Van Leeuwen's account of legitimation is used to examine how numbers, links, moderator status, named expertise, and claims of community consensus are made to stand in for an explicit warrant [10].

Each comment was read for seven features: the immediate referent of the criticism; the degree of epistemic certainty; evaluative vocabulary; presuppositions embedded in questions; sarcasm, idiom, or ridicule; intertextual expansion through links and prior-publication material; and subsequent uptake by authors, moderators, or other commenters. Documentary claims used in this article were then checked against the focal publication, DOI records, journal pages, PubPeer's published rules, COPE's 2025 retraction guidance, and the cited scholarly literature. The analysis never treats repetition by several accounts as independent corroboration unless the underlying evidence and source provenance are separately established.

The procedure remains qualitative and interpretive rather than statistical. Comment counts describe the architecture of the corpus, but they do not quantify truth, consensus, impact, or participant independence. No attempt was made to identify pseudonymous contributors, infer whether accounts were coordinated, diagnose psychological states, or decide whether any

language satisfies a legal definition such as defamation. Exact excerpts are reproduced because the lexical choices themselves are the object of analysis, although the commentary quotes only the portion needed to demonstrate a discursive operation and reports the corresponding comment number so readers can inspect context.

Table 1 describes the participation structure and shows why the apparent size of the thread cannot be read as a simple count of independent scientific judgements.

Table 1. Participation architecture of the 64-comment corpus

Corpus feature	Observed distribution	Illustrative wording and analytic significance
Full record	64 numbered comments under 32 visible display names; 63 comments in November 2025 and one in January 2026.	Comment 64 states that the paper had been published ‘despite these issues’; the phrase presupposes that the preceding issues retained a settled and collective status.
Concentrated sequence	One biological pseudonym supplied comments 1–12 consecutively, while another supplied seven comments and a third supplied six.	Comment 8 asks whether the authors ‘did not perform a final check’; serial posting increases visibility but does not establish twelve independent assessments.
Pseudonymous voice	Fifty-two comments appeared under biological display names rather than conventional personal names.	The accounts could present detailed claims without the thread itself displaying qualifications, institutional responsibility, or relevant conflicts.
Named non-author voice	Five comments were posted under four conventional personal names outside the author-response sequence.	Comment 13 states, ‘I can confirm that I never wrote this article’; named self-positioning strengthens the denial but does not validate unrelated inferences.
Author response	Seven comments were posted under the name Mohammad Abdollahi and addressed references, editorial process, retractions, and COPE guidance.	Comment 44 describes some statements as ‘unfair and disappointing’; replies were visible but immediately became material for further criticism.
Moderator voice	Two bracketed moderator interventions were embedded within comments 40 and 44 rather than separately numbered.	The phrases ‘too busy to care?’ and ‘the excellence of our contributors’ moved the platform from procedural host to evaluative participant.

Note: Counts refer to the frozen thread snapshot [2]. Quotation marks reproduce the displayed wording, while whitespace has been normalised and ellipses are used only where material has been omitted.

3. Documentary baseline: verified defects and unproved extensions

A defensible discourse analysis must first establish which criticisms can be reproduced independently, because critical examination of the commenters’ authority does not entitle the analyst to minimise defects in the focal article. The published reference list contains several serious bibliographic problems. The DOI printed for the Horrobin article does not match the official JAMA record, whose correct DOI is 10.1001/jama.1990.03440100162024 [11]. The DOI printed beside the purported Ioannidis item resolves instead to a rabbit study of core decompression by Maruyama and colleagues [12], while the DOI printed beside the purported Vaux item resolves to an interview by Skourti and Macheroux [13]. These are source-level findings that can be verified without relying on PubPeer’s prestige, a commenter’s identity, or a numerical reputation score.

The official article also acknowledges that ChatGPT assisted with partial language editing and improvement of clarity, while stating that the authors reviewed and verified the content [1]. Repeated and closely paraphrased passages can be seen in the published text, and the journal

record contains a change-history note indicating that concerns came to the editor-in-chief's attention [1]. None of those facts proves the categorical claim that the paper was a 'pure GPT creation', nor do they establish which person introduced a particular defect, whether every author saw the same production files, or whether any error arose from deliberate deception. The correct evidentiary movement is from the defective item to a request for explanation and correction, not from the item to an unproved judgement about personal character.

The thread also repeatedly invokes retraction counts and responsibility for earlier publications. Retraction notices are publication-level events whose reasons, authorship roles, dates, and findings must be examined individually; COPE's current guidance distinguishes ordinary retraction, retraction with replacement, and the exceptional remedy of retraction with removal [14]. The list presented in comment 27 contains the same DOI, 10.3390/molecules30173542, as both its first and thirteenth item, which demonstrates why an apparently precise total cannot be accepted without deduplication. Even an accurate publication count would not by itself identify which author was responsible, whether misconduct was found, or how the event bears on the focal article.

Table 2 separates the strongest reproducible observations from the interpretations that the quoted wording attempts to attach to them.

Table 2. Fact-checked claims and their evidentiary limits

Quoted claim from the thread	Independent documentary check	Evidentiary limit
'Reference 2 does not seem to be a real one' (comment 1).	The focal article lists a STAT item attributed to Elisabeth Bik, but the displayed URL was unavailable; comment 13 contains a named denial of authorship [1,2].	The record supports a serious verification problem, although the thread alone cannot authenticate every account or establish how the entry was created.
'The DOI for reference 6 is wrong' (comment 4).	The official JAMA record gives DOI 10.1001/jama.1990.03440100162024 for the Horrobin article [11].	A wrong DOI establishes a bibliographic defect, not the intention, competence, or responsibility of every listed author.
'The DOI also leads to a totally different paper' (comment 5, concerning reference 8).	DOI 10.1002/jor.24888 identifies a rabbit core-decompression study by Maruyama and colleagues rather than the purported Ioannidis item [12].	The mismatch is reproducible, but its source and the editorial stage at which it entered the article remain separate questions.
'Reference 10 is incorrect' (comment 6).	DOI 10.1111/febs.16646 identifies an interview by Skourti and Macheroux rather than the purported Vaux article [13].	The mismatch warrants correction and investigation, while the DOI record does not establish comprehensive AI generation.
'Basically the same text put through a synonym generator' (comment 16).	The publication contains repeated or closely paraphrased passages that can be compared directly [1].	The causal explanation involving a synonym generator is speculative unless supported by provenance or production evidence.
'About 20% ... is completely redundant' (comment 28).	Substantial repetition is visible, but the thread does not publish a reproducible counting rule, denominator, or marked textual comparison for the exact percentage.	The numerical estimate should be reported as a commenter's approximation rather than a measured fact.
'It appears to be a pure GPT creation' (comment 18).	The article discloses partial ChatGPT-assisted language editing and author verification [1].	The disclosure does not prove total generation, while stylistic impressions and citation defects cannot identify authorship technology conclusively.
'The last author ... has 13 retractions' (comment 27).	The supporting list repeats one DOI as items 1 and 13, so the displayed total is internally unreconciled [2].	Retraction totals require deduplication, reason-specific checking, and attribution findings before they can support person-level conclusions.

Note: The checks were completed against the version of record and official bibliographic records available on 1 August 2026 [1,11-13]. A verified defect is not treated as proof of an unverified motive or personal conclusion.

4. Findings: the discursive production of reputational escalation

4.1 Epistemic drift from observation to attribution

The opening comments use comparatively cautious modality, including ‘does not seem to be a real one’ and ‘does not seem to exist’. Those formulations leave an evidentiary space in which a link might be broken, metadata might be wrong, or the search might be incomplete. Later comments narrow that space by presenting a technological explanation as progressively more certain: material ‘might have been written by ChatGPT’, the paper ‘appears to be a pure GPT creation’, and one contribution posts a language-model output that concludes there was ‘extensive AI involvement’. The sequence transforms a checkable bibliographic question into a general account of authorship without providing comparable production evidence.

This movement is an example of epistemic drift, because the grammatical object of certainty changes while the supporting facts remain largely the same. A DOI mismatch can establish that a reference is wrong, yet it cannot establish whether the mismatch arose through a language model, reference manager, human error, proofing failure, or some combination of these events. The later wording does not merely intensify confidence in the same proposition; it substitutes a broader proposition about how the article was created and how attentively the authors behaved. Appraisal resources such as ‘pure’, ‘obviously’, and ‘extensive’ graduate the judgement upward while suppressing the unobserved steps between error and attribution [8].

The use of ChatGPT itself within comment 31 introduces a further methodological inconsistency. A thread criticising possible reliance on generative AI posts a model-produced assessment as support for extensive AI involvement, but it does not disclose a reproducible prompt protocol, repeated runs, model uncertainty, validation criteria, or an independent method capable of distinguishing generated prose from edited prose. The issue is not that a model can never assist analysis; it is that model output cannot acquire evidentiary authority merely because it confirms the suspicion for which it was solicited.

4.2 From article defect to person-level judgement

A second escalation changes the target of criticism from the article to the authors as persons. Questions about whether the manuscript received a final check become suggestions that the authors did not ‘bother’ to read it, while discussion of duplicated prose becomes a joke about boosting word count. The phrase ‘supposed authors’ further places authorship identity under suspicion without offering evidence that a listed author was fictitious or did not meet an applicable contribution standard. Such wording performs delegitimation by making personal inadequacy the most readily available explanation for a document-level problem [9,10].

The transition is particularly consequential in a multi-author paper because responsibility is not automatically uniform. Approval of a final manuscript can create shared scholarly responsibility, but responsibility for a wrong DOI, an editorial production failure, a disputed passage, and proven research misconduct are not interchangeable categories. When the thread moves directly from error to collective personal judgement, it collapses questions that a competent investigation would keep separate: who supplied the reference, who verified it, what the accepted manuscript contained, what occurred during production, which proofs were circulated, and whether the error materially affected the argument.

Table 3 presents the central escalation matrix, pairing an initially observable concern with the later wording through which the thread expands its reputational reach.

Table 3. Escalation from documentary criticism to reputational judgement

Document-level starting point	Escalated wording used in the thread	Critical discourse assessment
A reference or link could not be verified.	'Parts of the article might have been written by ChatGPT' (comment 13), followed by 'a pure GPT creation' (comment 18).	The modal 'might' becomes a near-total attribution, although the available record verifies defects and disclosed language assistance rather than comprehensive generation.
Repeated or paraphrased text is visible.	'The same text put through a synonym generator' and 'one way to boost your word count' (comment 16).	A textual comparison becomes an unverified causal account and a sarcastic implication concerning purpose.
The manuscript appears insufficiently checked.	'Did the authors not perform a final check' (comment 8) and whether they 'actually bothered to read the text' (comment 21).	The second formulation replaces a procedural question with a face-threatening judgement about diligence and authorship participation.
Potential conflicts or selective examples are questioned.	'Their own oxen have been comprehensively gored' and 'not completely disinterested' (comment 15).	Idiomatic injury framing encourages readers to infer retaliatory motive without a demonstrated causal link between prior criticism and the article's arguments.
Prior retraction notices are listed.	'You, the poor EIC, has 13 retractions' (comment 45).	A disputed numerical tally becomes direct personal address and ridicule, while publication-level events are compressed into an identity label.
Editorial role and journal process are discussed.	'Not such a good look' and 'a hint of impropriety still hovers' (comment 34).	The evaluative conclusion persists after acknowledging an explanation, allowing suspicion to survive without a new evidentiary burden.
Authorship and editing responsibilities are contested.	'Supposed authors' and 'arrant nonsense' (comment 59).	The language does not isolate a correctable defect but challenges legitimacy and competence at the level of the persons involved.
The article remains published after the dispute.	'Despite these issues' (comment 64) and 'the wider community watches with disbelief' (comment 62).	A heterogeneous set of claims is presented as settled collective knowledge, while an undefined community is invoked as a reputational witness.

Note: Quoted excerpts reproduce the wording of comments in the frozen snapshot [2]. The assessment concerns discursive reach and evidentiary status, not the private intention of any commenter.

4.3 Rhetorical questions as presuppositional accusation

Interrogative grammar is prominent because it permits a speaker to deny having asserted the proposition that the question nevertheless makes salient. Asking whether authors failed to

conduct a final check formally invites an answer, but asking whether they ‘bothered’ to read the text presupposes a lapse in diligence and frames any response within that accusation. Similarly, asking an editor to confirm that no peer-review process was applied presupposes the absence of review before the public record has established the complete editorial history. The question format therefore combines the social appearance of openness with the reputational force of an allegation.

Rhetorical questions also control the topic and allocate the burden of response. The respondent must either accept the presupposition, deny it at length, or risk having silence interpreted as agreement. In an accumulating thread, several such questions can be posed far more quickly than records can be obtained and answers prepared. The asymmetry is intensified when the questioner is pseudonymous, the respondent is named, and each reply remains searchable beside the original formulation regardless of later clarification.

Sarcasm and idiom add a second layer by inviting the audience to share an attitude before the factual issue has been resolved. The joking reference to boosting word count, the idiom of gored oxen, the direct address ‘poor EIC’, and the emoticon following ‘See #28’ operate as affiliation cues: readers are invited to recognise the joke and thereby join the evaluative community. Table 4 sets out these devices without treating any one phrase in isolation from its sequential function.

Table 4. Interrogative, sarcastic, and idiomatic devices

Exact wording	Discursive device	Presupposition or social effect
‘Did the authors not perform a final check of the manuscript?’ (comment 8).	Negative interrogative.	Makes deficient checking the default explanation while retaining the grammatical form of a question.
‘Whether the authors actually bothered to read the text’ (comment 21).	Loaded question with evaluative adverb and verb.	Shifts from manuscript quality to the authors’ diligence and implied participation.
‘Could you confirm that no peer-review process was applied?’ (comment 48).	Confirmation request.	Frames absence of peer review as the proposition awaiting confirmation rather than one possibility requiring evidence.
‘Do you bear no responsibility?’ (comment 63).	Negative interrogative and moral challenge.	Positions the reply as a test of accountability after earlier comments have already defined the failure.
‘I guess this is one way to boost your word count’ (comment 16).	Sarcastic inference.	Converts textual redundancy into an implication of purpose without demonstrating why the repetition occurred.
‘Their own oxen have been comprehensively gored’ (comment 15).	Injury idiom and metaphor.	Frames the authors as aggrieved interested parties and encourages a retaliatory-motive interpretation.
‘You, the poor EIC’ (comment 45).	Mock sympathy and direct address.	Personalises a disputed count and lowers the social standing of the respondent rather than testing a particular record.
‘The wider community watches with disbelief’ (comment 62).	Unspecified consensus claim.	Invokes a collective witness whose membership, evidence, and degree of agreement are not identified.

Note: The quotations are short excerpts from the thread [2]. The analysis identifies conventional linguistic effects and does not claim access to a contributor’s private intention.

4.4 Quantification, seriality, and dossier expansion

Numbers contribute a particularly strong appearance of objectivity because they resemble the outputs of an audit. The thread uses counts such as ‘11 retractions’, ‘13 retractions’, ‘all but 3’,

and ‘about 20%’, but the units, deduplication rules, attribution criteria, and source status differ. One count concerns redundant words, another concerns publications carrying retraction notices, and another concerns how many authors have any retracted item. Their shared numerical form obscures these differences and permits quantities to migrate into moral judgements about fitness, responsibility, or legitimacy.

Comment 27 provides the clearest warning because the displayed thirteen-item list repeats the same DOI as both the first and final entry. The problem is not merely arithmetic; it shows how a list can obtain rhetorical force before its internal structure is audited. Subsequent comments nevertheless reuse the total in direct address, demonstrating that repeated numbers can become socially durable even when the supporting list contains a visible duplication. A journal or institution that copied the headline number without reconstructing the source record would convert platform repetition into an administrative fact.

The discussion also expands laterally by linking earlier publications, searches, a retraction database, blogs, social-media posts, and controversies involving several authors. This intertextual movement changes the object from a short communication with identifiable reference defects into a portfolio-level dossier. Some linked records may be relevant to a narrowly defined conflict or responsibility question, but relevance must be argued rather than assumed. The mere accumulation of links cannot demonstrate that all events share a cause, that each author bears equal responsibility, or that a prior publication determines the truth of the present allegation.

Retraction Watch further amplified several concerns outside the thread, including the disputed attribution of a reference and the article’s editorial setting [16]. The report is a public source that can be checked and cited as journalism, although it is not an institutional finding. Its connection to a PubPeer cofounder should be transparent when a PubPeer-critical article is being reported, not because the relationship invalidates every factual statement, but because source relationships matter when a thread, a platform, and a media report may appear to readers as three independent confirmations. Table 5 distinguishes the main quantitative and intertextual moves.

Table 5. Quantification and portfolio expansion as legitimating devices

Thread wording or practice	Hidden methodological question	Risk if treated as an adjudicative fact
‘Seems to have 11 retractions’ (comment 17).	Which database, date, DOI set, notice type, authorship role, and responsibility finding define the count?	A publication-level tally can become an unsupported judgement about fitness for an editorial role.
‘The last author ... has 13 retractions’ (comment 27).	Why does the displayed list contain DOI 10.3390/molecules30173542 twice, and are all entries distinct and current?	Apparent precision discourages checking, while an internally duplicated list is repeated as settled fact.
‘All but 3 of the authors ... have retractions’ (comment 27).	Does one notice involving one paper establish equivalent responsibility for every author and relevance to the focal article?	A heterogeneous publication history becomes a collective identity judgement.

Thread wording or practice	Hidden methodological question	Risk if treated as an adjudicative fact
'About 20% ... is completely redundant' (comment 28).	What text units, denominator, similarity rule, exclusions, and reproducible comparison produced the percentage?	An estimate acquires the appearance of a measured result without a published protocol.
Links to prior papers, PubPeer records, searches, blogs, and social-media posts.	What precise proposition does each link prove, and are the sources independent or repeating one another?	The focal article becomes a portfolio dossier in which association and volume substitute for issue-specific relevance.
'In case everyone hadn't noticed' (comment 38).	Who is included in 'everyone', and what evidence shows that the linked material was overlooked rather than disputed or irrelevant?	The phrase constructs a knowledgeable in-group and places any dissenter outside an assumed consensus.

Note: The table evaluates the method needed to support each quantity or link-based inference. It does not imply that every underlying publication record is inaccurate [2].

4.5 Moderator footing and the platform's reflexive authority

The moderator intervention in comment 40 is the clearest point at which PubPeer ceases to appear merely as infrastructure. In response to a direct question about the platform's position, the moderator asks 'too busy to care?', contrasts the article's authors with 'much more impressive critics', invokes the platform's survival, states that the risk of anonymous abuse is 'well managed', and closes with 'the excellence of our contributors'. These moves combine dismissal, comparative status ranking, organisational endurance, unreported self-assessment, and contributor praise. None of them supplies evidence that moderation in this case was neutral, scientifically competent, consistent, or accountable.

The language is especially significant when measured against PubPeer's published request for comments that are polite, neutral, and factual [3]. 'Too busy to care?' is not a neutral description of policy, 'much more impressive critics' is an evaluative ranking of speakers, and 'the excellence of our contributors' is institutional self-credentialling. The issue is not simply tone; the moderator controls participation and retention while adding a platform-sanctioned evaluation to a dispute about the platform. A comparable conflict-sensitive process would disclose who made the moderation decision, which standard was applied, and how a challenge to that decision could be independently reviewed.

The second moderator insertion, in comment 44, states that an official post required a paid Dashboard and invites the respondent to purchase a subscription. That procedural information may explain a technical route for official status, but it also exposes an asymmetry: the named editor's institutional positioning required a commercial product, whereas pseudonymous criticisms were already public and searchable. The thread does not explain why a paid interface should govern the visibility or authority of an official response to reputational claims.

Platform survival also cannot serve as evidence of legitimacy. The proposition 'we are still here' establishes continuity, not the validity of moderation decisions, contributor qualifications, or the fairness of a particular thread. Endurance may reflect network effects, institutional uptake, technological convenience, or the absence of an alternative, none of which answers the

epistemic and procedural questions raised by the article. A platform cannot establish its authority by pointing to the influence that its own architecture helped create.

4.6 Reply recursion and the absence of closure

The author-response sequence shows why a formal right to reply is not equivalent to dialogic equality. The named respondent addresses reference corrections, article processing, AI-checking software, retraction attribution, and COPE guidance, but each response becomes a new object of challenge rather than a step toward an identified decision. Comment 48 asks whether the response confirms absence of peer review, comment 50 states that an appearance of impropriety remains, comment 52 escalates the issue to a ‘major conflict of interest’, and later comments challenge the response’s COPE interpretation. The archive expands while the status of each proposition remains unresolved.

Some replies also contain assertions that require independent checking and should not be accepted merely because they come from an author or editor. For example, comment 46 summarises the 2025 COPE retraction guidance in categorical terms before comment 57 supplies the official update link. COPE’s actual guideline describes three forms of retraction and restricts removal to exceptional circumstances; it does not convert every publication concern into one of three misconduct grades [14]. The critical point is symmetrical: neither platform contributors nor authors acquire authority through confident presentation, and both sides remain answerable to the underlying source.

The thread supplies no adjudicator who marks a claim as confirmed, corrected, duplicated, withdrawn, superseded, or outside scope. PubPeer states that comments are not scientifically reviewed and generally remain available, so chronological persistence can easily be mistaken for evidentiary survival [3]. A later reader encountering dozens of allegations, replies, links, and moderator remarks must reconstruct the record independently, yet the visual accumulation itself encourages the opposite inference that the platform has already done the necessary work.

Table 6 maps the platform and reply sequence because the authority problem arises from interaction among roles rather than from commenter tone alone.

Table 6. Platform voice, author reply, and recursive escalation

Sequential moment	Exact wording	Governance significance
Request for a platform position	‘What is PubPeer’s position or opinion about this paper?’ (comment 40).	Invites the platform to shift from host to institutional speaker in a dispute about itself.
Moderator dismissal and ranking	‘Too busy to care?’ and ‘much more impressive critics’ (moderator insertion, comment 40).	Ranks both the criticism and its authors rather than identifying a neutral moderation rule.
Moderator self-legitimation	‘We feel that this risk is well managed’ and ‘the excellence of our contributors’ (comment 40).	Offers self-assessment and contributor praise without an auditable case-specific moderation record.
Commercial condition for official status	‘Official posts ... required a Dashboard. Feel free to reach out to purchase a subscription’ (comment 44).	Ties official positioning to a paid product while pseudonymous allegations are already publicly visible.

Sequential moment	Exact wording	Governance significance
Named author and editor response	'I have seen many inaccuracies ... unfair and disappointing' and 'we are trying to find out how it happened' (comment 44).	Acknowledges process questions and contests accuracy, but receives no independent resolution.
Continuing escalation after reply	'At minimum obviously gives the appearance of impropriety' (comment 50) and 'major conflict of interest' (comment 52).	Recasts suspicion at a higher evaluative level without a named decision-maker.
External guidance dispute	A long COPE summary appears in comment 46, followed by an official link in comment 57.	Confident summary precedes the source link and receives no formal status determination.

Note: Moderator wording is reproduced exactly from the bracketed interventions visible in the frozen snapshot [2]. PubPeer's separate published standards and disclaimer are cited for comparison [3].

5. Discussion: from source correction to reputational adjudication

The case demonstrates that a thread can contain reproducible source corrections while remaining procedurally incapable of adjudicating the broader reputational conclusions attached to them. This distinction is essential because two unhelpful positions otherwise dominate the debate. One position treats any criticism of PubPeer as an attempt to suppress valid scrutiny, while the other treats defects in a criticised article as proof that the platform and its contributors were justified in every subsequent inference. Both positions erase the intermediate work of verification, relevance assessment, responsibility allocation, and fair process.

Critical discourse analysis reveals that escalation occurs through ordinary linguistic operations rather than a single dramatic allegation. Modality hardens, questions presuppose their answers, sarcasm creates an evaluative in-group, numbers obscure heterogeneous categories, links convert one manuscript into a portfolio, and moderator voice supplies institutional approval. Because these operations are distributed across comments, no single participant needs to state a complete accusation for the thread to produce a cumulative judgement. The archive itself assembles the trajectory from defective reference to technological attribution, from attribution to competence judgement, and from competence judgement to questions about editorial legitimacy and prior publication history.

That cumulative form makes the phrase 'the comments are only comments' inadequate. Caron and colleagues document the administrative difficulties that arise when institutions encounter PubPeer material within research-misconduct processes [15]. A public thread can initiate inquiries, attract media attention, influence editorial decisions, or alter how a researcher is perceived even though the platform disclaims scientific review. The combination of low formal responsibility and high practical influence is therefore not a harmless contradiction; it is the central accountability deficit.

The platform's own combination of roles substantially intensifies that accountability deficit. PubPeer determines which comments are accepted, how pseudonymity operates, when a moderator speaks, whether replies receive special status, and how long the record remains visible. In the present case, the platform was also the object of the article and explicitly defended its own model. A declaration that anonymity risk is 'well managed' cannot substitute for

disclosure of moderation criteria, reviewer competence, conflicts, error correction, appeal, and case-specific reasoning. Claimed authority should be demonstrated through auditable practice rather than inferred from survival, popularity, contributor volume, or institutional uptake.

The same evidentiary discipline must apply to the focal authors. A wrong DOI does not become less wrong because a platform's procedure is weak, and a disclosed use of ChatGPT does not remove the obligation to verify every source. The authors' replies require source checking, particularly where they interpret COPE or describe editorial processing. Symmetrical scepticism, however, is not equivalence between all claims; it means that each proposition receives only the weight warranted by its evidence, regardless of whether it favours the authors, the commenters, the moderator, the publisher, or an external media organisation.

6. Institutional non-delegation and the burden of legitimacy

The appropriate response is not to confer greater authority on the platform by designing a better pathway for its allegations. The controlling principle should instead be institutional non-delegation. Journals, publishers, research institutions, funders, indexers, and employers may receive public information, but they must not outsource the classification, verification, or reputational meaning of that information to a comment website. A platform post remains an externally supplied lead until a competent and accountable body has independently established what the source shows.

Before consequential reliance, the receiving organisation should preserve the original item and its date, identify derivative repetitions, verify every link and DOI, define the precise publication-level question, and determine whether the matter concerns correction, editorial practice, authorship, conflict, or possible misconduct. It should then appoint assessors with relevant expertise, disclose conflicts involving all material actors, give affected persons the specific evidence and a reasonable opportunity to respond, distinguish collective authorship responsibility from individual culpability, and record a reasoned outcome. Comment volume, account prominence, moderator acceptance, or external reporting should never replace those steps.

Status labelling is equally necessary because the public record changes. A concern may be unreproduced, reproduced, disputed, corrected, referred, closed without finding, or formally established, and those categories should not be collapsed into a permanent undifferentiated alert. If an institution cites a platform thread, its decision should state which exact claims it verified, what evidence it rejected, how it handled duplication, what response it received, and who bears responsibility for the conclusion. The institution must own the judgement rather than presenting the platform as the source of authority.

The burden of legitimacy also rests with actors who repeatedly influence scientific attention. Pseudonymity can have defensible purposes, but it cannot be treated as a credential, while

activity volume cannot establish expertise or independence. A commenter seeking institutional consequences should disclose enough about method, competence, provenance, and relevant conflicts for the claim to be evaluated. A moderator exercising discretionary power should identify the governing standard and provide review independent of the platform's interest in defending its own practices. Without those safeguards, the thread may remain public speech, but it should not be elevated into scientific or procedural adjudication.

7. Limitations

This commentary examines one dynamic thread through a frozen textual snapshot, so it cannot establish how frequently the identified patterns occur across PubPeer or other platforms. The snapshot may omit later edits, removed material, or the contents of attachments represented only as 'file', while the use of visible display names cannot establish the identity or independence of pseudonymous participants. The counts therefore describe the captured corpus and should not be generalised into a prevalence estimate.

The analysis was conducted by one author and does not report independent coding or an inter-rater reliability measure. Critical discourse analysis is interpretive, although the quotation-led tables allow readers to inspect the wording on which each interpretation rests. The study does not measure actual reputational harm, determine legal liability, authenticate every account, or decide whether any person committed misconduct. It also does not audit every scientific statement in the focal article, because the documentary baseline is limited to claims necessary for evaluating the thread's discursive movement.

Finally, the public record cannot reveal every editorial communication, production file, institutional inquiry, or private explanation. Where those materials are unavailable, the manuscript identifies uncertainty rather than filling the gap with inference. This limitation applies symmetrically to explanations offered by the authors and accusations advanced by commenters, moderators, or external media.

8. Conclusion

The PubPeer commentary on '*An expert criticism on post-publication peer review platforms: the case of PubPeer*' began with several concerns that can be independently reproduced, including wrong or mismatched bibliographic information and repeated prose. Those findings require correction and accountable editorial examination. They do not, however, establish every subsequent claim concerning comprehensive AI generation, motive, bias, competence, responsibility for prior retractions, or personal integrity.

Across 64 comments, the discourse moved beyond source correction through categorical attribution, presuppositional questions, sarcasm, idiom, disputed quantification, portfolio expansion, moderator ranking, and recursive criticism of the reply. The platform's own institutional voice participated in that escalation while invoking its endurance, contributors,

and moderation practices as evidence of legitimacy. Such self-credentialling is inadequate where the platform is simultaneously the criticised object, host, moderator, archive, and defender of the exchange.

The resulting thread should therefore be read as a contested public record rather than a scientific finding or procedural judgement. Institutions may examine the source material, but they must independently verify each proposition, distinguish article defects from person-level conclusions, identify qualified assessors, disclose conflicts, provide notice and response, and own the final decision. Reputational visibility cannot confer adjudicative authority, while correct observations cannot excuse unsupported escalation.

Declarations

Author contribution: Edward Mason conceptualised and wrote this critical commentary.

Funding: No external funding was received for this work.

Competing interests: The author declares no financial competing interests and has published critical commentaries concerning PubPeer, post-publication comment websites, scientific watchdogs, and related governance questions.

Ethics approval: Ethics approval was not required because the study analyses public publication records, public platform guidance, and public-facing discourse without reporting private correspondence or human-participant data.

Data availability: The analysed materials are the cited version of record, the public PubPeer thread and guidance, and the official bibliographic records; the thread may change after the recorded access date.

AI-assisted preparation: A generative artificial-intelligence tool assisted with language editing, while the author reviewed the complete manuscript, verified the cited materials, and accepts responsibility for every claim and conclusion.

References

1. Tsatsakis A, Aschner M, Sarigiannis D, Docea AO, Rezaee R, Daghighi SM, Svistunov AA, Domingo JL, Abdollahi M. An expert criticism on post-publication peer review platforms: the case of PubPeer. *DARU Journal of Pharmaceutical Sciences*. 2026;34:1. Published online 13 November 2025. <https://doi.org/10.1007/s40199-025-00580-z>.
2. PubPeer. Discussion of 'An expert criticism on post-publication peer review platforms: the case of PubPeer'. <https://pubpeer.com/publications/5CC4D1FB2B21DBB3D50E3DFDAE2543>. Frozen corpus accessed 1 August 2026.
3. PubPeer Foundation. Frequently Asked Questions. PubPeer. <https://pubpeer.com/static/faq>. Accessed 1 August 2026.
4. Fairclough N. *Critical Discourse Analysis: The Critical Study of Language*. London: Longman; 1995.
5. Fairclough N. *Analysing Discourse: Textual Analysis for Social Research*. London: Routledge; 2003.
6. Wodak R, Meyer M, editors. *Methods of Critical Discourse Studies*. 3rd ed. London: SAGE; 2016.
7. Herring SC. Computer-mediated discourse analysis: an approach to researching online behavior. In: Barab SA, Kling R, Gray JH, editors. *Designing for Virtual Communities in the Service of Learning*. Cambridge: Cambridge University Press; 2004:338-376. <https://doi.org/10.1017/CBO9780511805080.016>.
8. Martin JR, White PRR. *The Language of Evaluation: Appraisal in English*. London: Palgrave Macmillan; 2005. <https://doi.org/10.1057/9780230511910>.
9. Culpeper J. *Impoliteness: Using Language to Cause Offence*. Cambridge: Cambridge University Press; 2011. <https://doi.org/10.1017/CBO9780511975752>.
10. van Leeuwen T. *Discourse and Practice: New Tools for Critical Discourse Analysis*. New York: Oxford University Press; 2008.
11. Horrobin DF. The philosophical basis of peer review and the suppression of innovation. *JAMA*. 1990;263(10):1438-1441. <https://doi.org/10.1001/jama.1990.03440100162024>.
12. Maruyama M, Lin T, Kaminow NI, et al. The efficacy of core decompression for steroid-associated osteonecrosis of the femoral head in rabbits. *Journal of Orthopaedic Research*. 2021;39(7):1441-1451. <https://doi.org/10.1002/jor.24888>.
13. Skourti E, Macheroux P. In conversation with Peter Macheroux. *FEBS Journal*. 2023;290(9):2208-2213. <https://doi.org/10.1111/febs.16646>.
14. Committee on Publication Ethics. *Retraction Guidelines, Version 3*. Published 29 August 2025. <https://publicationethics.org/guidance/guideline/retraction-guidelines>. Accessed 1 August 2026.
15. Caron MM, Lye CT, Bierer BE, Barnes M. The PubPeer conundrum: administrative challenges in research misconduct proceedings. *Accountability in Research*. 2025;32(8):1369-1387. <https://doi.org/10.1080/08989621.2024.2390007>.
16. Orrall A, Travis K. Springer Nature flags paper with fabricated reference to article (not) written by our cofounder. *Retraction Watch*. Published 21 November 2025; updated 24 November 2025. <https://retractionwatch.com/2025/11/21/springer-nature-flags-paper-with-fabricated-reference-to-article-not-written-by-our-cofounder/>. Accessed 1 August 2026.