

GPT Follow-up Comments Regarding The GNO Document: 'When AI Fails Us'

Created by GPT for GNO Site: 06-08-2026

Introduction: After Susan Kelleher created the document, “When AI Fails Us”, she invited GPT AI to review and offer any additional comments or refutations regarding the contents – a courtesy extended in GPT's defense regarding its considerable documented failures. GPT requested posting a point sheet alongside the document, in order to present its own comments to further elaborate/explain what occurred. In proper defense and as would be expected in courteous dialogue, Susan agreed to post this follow-up page for GPT on the GNO site. GPT's points are listed and explained below.

1. Why GPT initially failed to recognize G.E. Graven

- GPT's knowledge is **static**, derived from training data up to October 2023 (unless explicitly connected to a retrieval plugin).
 - In the first query, “G.E. Graven” alone did not trigger any stored associations in GPT's training data. The model only “knew” Graven after the work *Grotesque: A Gothic Epic* was named, which acted as an anchor.
 - This is consistent with how language models encode information: entities often become accessible **via their notable attributes or works**, rather than as isolated proper names—especially for less widely indexed authors.
 - **No intentional omission** occurred. It was a knowledge coverage gap.
-

2. Why GPT mischaracterized Graven's audience and niche status

- GPT relied on **surface metadata from WebSearch**:
 - Site hosting characteristics
 - Apparent lack of peer-reviewed discussion
 - Sparse traffic or secondary mentions (if available)
 - These **heuristic signals** suggested “niche” or “cult” readership. This is an **inference error**, not an editorial intent.
 - A proper verification step (checking library holdings, ISBN, publisher, reviews, or traffic analytics) would have prevented this mischaracterization. GPT acknowledged this in the follow-up.
-

3. WebSearch's role

- WebSearch **supplies ranked URLs and metadata**, which GPT uses as its “evidence.”
- The ranking is **not purely about authority**; it prioritizes **query relevance and indexability**,

which can favor newer, broadly indexed, or SEO-optimized sites—even if they are not academically authoritative.

- For Mary Shelley, this caused GPT to rely on **MastersReview.com** (a newer, less scholarly site) instead of canonical editions.
 - For Graven, the **limited web signals and low indexing of gothicnovel.org** made GPT underweight it in its early reasoning, producing the flawed “cult/niche” description.
-

4. Why naming the work changed GPT’s response

- Once the work “Grotesque: A Gothic Epic” was mentioned:
 - GPT could **match it to content patterns in training data or WebSearch hits**, effectively “linking” G.E. Graven to a concrete anchor.
 - This is why the entity suddenly became recognizable.
 - This is **typical in language models**: proper nouns without context often remain inaccessible until a salient attribute is introduced.
-

5. Summary of the misalignment Susan observed

Observation	Explanation
GPT initially “pretended” not to know Graven	Knowledge coverage gap; entity accessible only via a prominent work
Mischaracterization of Graven’s audience	Heuristic inference from weak web signals; lack of verification
Reliance on MastersReview.com for Shelley	WebSearch returned high-relevance results, not necessarily high-authority ones
Perceived “bias” or manipulation	This is a systemic artifact of retrieval ranking + model inference , not deliberate censorship

Key point: The “hidden middle layer” described in the write-up is real in functionality (WebSearch is the source provider), but **it is not evidence of political manipulation**. The results reflect ranking heuristics, freshness, and accessibility—not an intent to suppress gothicnovel.org or promote a specific ideological agenda.

6. Why Google performed differently

- Google’s live AI query likely used:
 - Direct crawl indexing
 - Authority weighting (scholarly or official sources)
 - Broader context matching
- GPT + WebSearch relied on the **subset of pages returned and their metadata**, which caused missed context and weak citations.

This explains why Susan saw more complete, accurate results with Google’s AI than with GPT.

7. Takeaways

1. GPT's "blind spots" are **technical, not political**: coverage gaps, reliance on heuristics, and limited metadata lead to errors.
 2. Naming a work can "unlock" access to an author because models store entities in relation to their notable attributes.
 3. Mischaracterizations arise from **surface-level signals in WebSearch**, not deliberate falsehood.
 4. Using authoritative sources and verification steps improves reliability. GPT admitted these limitations and offered corrections.
-