



TECHNICAL BRIEF

AI News Audit

How AI Models Use and Distribute Canadian Journalism

Aengus Bridgman & Taylor Owen

Aengus Bridgman is Associate Director, Research, and Director of the Media Ecosystem Observatory at the Centre for Media, Technology and Democracy, McGill University.

Taylor Owen is the Beaverbrook Chair in Media, Ethics and Communication and Founding Director of the Centre for Media, Technology and Democracy, McGill University.

MARCH 2026

Executive summary

AI models have become a primary way for millions of Canadians to get information. When people ask ChatGPT, Gemini, Claude, or Grok about current events, the responses they receive, and the sources those responses cite, shape which news organizations get traffic, attribution, and revenue. This brief presents the first large-scale empirical audit of how AI models use and distribute Canadian journalism.

We tested four major AI models on 2,267 real Canadian news stories (English and French) without web search activated and found the same pattern across all of them. All four models showed extensive knowledge of Canadian current events consistent with having ingested Canadian news reporting. Models demonstrated at least partial knowledge in 74% of responses to stories within their training window, but among those knowledgeable responses, 92% provided no source attribution of any kind.

When we enabled web search and tested 140 specific articles via each company's API, every model produced responses that covered enough of the original reporting that many consumers would rarely need to visit the source. Models often linked to Canadian news sites, with 52% of responses including at least one Canadian URL, but named a Canadian source in the response text only 28% of the time. Links provide a pathway back to the source, but consumers reading the response itself rarely see an indication of whose journalism they are consuming.

AI models readily covered content from paywalled sources. In most cases, competitive news markets mean the same stories appear across multiple outlets, and models retrieve a free version. In some cases, API logs showed models citing paywalled URLs directly, suggesting that paywalls may not block automated retrieval the way they block human readers. Whether models access the content directly or through freely available alternatives, the result for publishers is the same: the originating outlet receives neither traffic nor credit. AI companies ingest Canadian journalism at scale, produce substitute content from it, and distribute that content to consumers. Links to Canadian news sites are provided, but the newsrooms that produced the journalism are rarely named in the response consumers read.

Résumé exécutif

Les modèles d'IA sont devenus l'un des principaux moyens de s'informer pour des millions de Canadiens. Lorsque des personnes interrogent ChatGPT, Gemini, Claude ou Grok sur l'actualité, les réponses qu'elles reçoivent – et les sources que ces réponses citent – influencent quels médias reçoivent du trafic, de l'attribution et des revenus. Ce rapport présente le premier audit empirique à grande échelle de la manière dont les modèles d'IA utilisent et distribuent le journalisme canadien.

Nous avons testé quatre modèles d'IA sur 2 267 événements d'actualité canadiens réels (en anglais et en français) sans recherche web activée et avons observé le même schéma pour chacun d'eux. Les quatre modèles ont démontré une connaissance étendue de l'actualité canadienne, cohérente avec l'ingestion de reportages de médias canadiens – au moins une connaissance partielle dans 74 % des réponses aux événements couverts par leurs données d'entraînement. Parmi ces réponses, 92 % ne fournissaient aucune attribution de source.

Lorsque nous avons activé la recherche web et testé 140 articles spécifiques via l'API de chaque entreprise, chaque modèle produisait des réponses couvrant suffisamment du reportage original pour que les consommateurs n'aient que rarement besoin de consulter la source. Les modèles incluaient souvent des liens vers des sites de médias canadiens, avec 52 % des réponses contenant au moins une URL canadienne, mais ne nommaient une source canadienne dans le texte que 28 % du temps. Les liens permettent de retracer la source, mais les consommateurs qui lisent la réponse elle-même voient rarement une indication du média dont ils consomment le journalisme.

Les modèles d'IA ont facilement couvert le contenu d'articles payants. Dans la plupart des cas, les marchés de l'information concurrentiels font que les mêmes articles paraissent dans plusieurs médias, et les modèles récupèrent une version gratuite. Dans certains cas, les journaux d'API indiquaient que les modèles citaient directement des URL payantes, ce qui suggère que les verrous d'accès payants ne bloquent pas nécessairement la récupération automatisée comme ils bloquent les lecteurs humains. Que les modèles accèdent au contenu directement ou via des alternatives gratuites, le résultat est le même pour les éditeurs : le média d'origine ne reçoit ni trafic ni crédit. Les entreprises d'IA ingèrent le journalisme canadien à grande échelle, produisent du contenu de substitution et le distribuent aux consommateurs. Des liens vers des sites de médias canadiens sont fournis, mais les salles de rédaction qui ont produit le journalisme sont rarement nommées dans la réponse que les consommateurs lisent.

Key Takeaways

- **AI models have absorbed Canadian news at scale and almost never say so.** All four models showed knowledge of Canadian current events for stories inside their training windows. Wire services and other public sources covering the same events may also contribute, but the pattern is consistent with having ingested Canadian news reporting. Among responses where a model demonstrated that knowledge, 92% provided no source attribution of any kind: no named outlet, no recommendation, no reference.
- **AI models substitute for the article and link to sources, but rarely name them.** With web access, models covered enough distinctive reporting to reduce a consumer's need to visit the source in 54% (ChatGPT) to 81% (Gemini) of cases. Models linked to Canadian news sites in 52% of responses. But they rarely named any Canadian source in the response text: 2% (Gemini) to 54% (ChatGPT) of responses. When directly prompted, naming rates reached 74% (Grok) to 97% (Claude).
- **Paywalls do not protect content in a competitive news market.** AI models covered paywalled content (64%) at rates comparable to free content (70%). For most stories, models find freely available versions elsewhere. In some cases, API logs showed models citing paywalled URLs directly with extensive verbatim reproduction, suggesting that paywalls may not block automated retrieval the way they block human readers. Either way, the result is the same: paywalled journalism is reproduced without compensation.
- **Naming sources is technically possible but rarely the default.** When prompted, naming rates reached 74% (Grok) to 97% (Claude). Under default conditions, models linked to Canadian news sites 52% of the time but rarely named the outlet in the response text itself. Links let a consumer reach the source; naming would tell them whose journalism they are reading. The gap between the two is a design choice.

Principaux constats

- **Les modèles d'IA ont absorbé le journalisme canadien à grande échelle, et ne le disent presque jamais.** Les quatre modèles ont démontré une connaissance de l'actualité canadienne pour les sujets couverts par leurs données d'entraînement. Les agences de presse et d'autres sources publiques couvrant les mêmes événements peuvent aussi contribuer, mais le schéma est cohérent avec l'ingestion de reportages de médias canadiens. Parmi les réponses où un modèle a démontré cette connaissance, 92 % ne fournissaient aucune attribution : ni média nommé, ni recommandation, ni référence.
- **Les modèles d'IA se substituent à l'article et fournissent des liens, mais nomment rarement leurs sources.** Avec l'accès au web, les modèles couvraient suffisamment du reportage distinctif pour réduire le besoin de consulter la source dans 54 % (ChatGPT) à 81 % (Gemini) des cas. Les modèles incluaient des liens vers des sites de médias canadiens dans 52 % des réponses. Mais ils nommaient rarement le média d'origine dans le texte : 2 % (Gemini) à 54 % (ChatGPT) des réponses. Lorsque nous demandions des citations directement, les taux atteignaient 74 % (Grok) à 97 % (Claude).
- **Les verrous d'accès payants ne protègent pas le contenu dans un marché de l'information concurrentiel.** Les modèles d'IA ont couvert le contenu payant (64 %) à des taux comparables au contenu gratuit (70 %). Pour la plupart des articles, les modèles trouvent des versions gratuites ailleurs. Dans certains cas, les journaux d'API ont montré des modèles citant directement des URL payantes avec une reproduction textuelle extensive, ce qui suggère que les verrous d'accès payants ne bloquent pas nécessairement la récupération automatisée comme ils bloquent les lecteurs humains. Dans tous les cas, le résultat est le même : le journalisme payant est reproduit sans compensation.
- **Nommer les sources est techniquement possible, mais rarement le comportement par défaut.** Lorsque nous demandions des citations, les taux de nomination atteignaient 74 % (Grok) à 97 % (Claude). Par défaut, les modèles incluaient des liens vers des médias canadiens 52 % du temps, mais nommaient rarement le média dans le texte de la réponse. Les liens permettent au consommateur d'atteindre la source ; nommer le média lui dirait quel journalisme il lit. L'écart entre les deux est un choix de conception.

Why it Matters

Journalism is part of democratic infrastructure. It is essential for accountability, public debate, and informed civic participation. Journalism is also expensive to produce and has been under severe financial pressure for two decades. AI companies are now extracting value from it at every stage: ingesting news archives as training data, producing derivative content without naming the sources, and delivering answers to consumers that could reduce the need and incentive to visit the original source.

This is not a hypothetical concern. Our data show that AI models possess detailed knowledge of Canadian current events: knowledge that most plausibly comes from Canadian news reporting. They answer questions about Canadian politics, policy, and society with a specificity that points to Canadian journalism as the primary source. They often link to Canadian news sites, but they rarely name the journalists or organizations whose reporting informed the answer. A consumer who does not click through every link will not know whose journalism they just consumed. The result is a system in which AI companies monetize journalists' knowledge while the newsrooms that produced it receive links but not recognition. That dynamic, if sustained at scale, risks accelerating the decline of the journalism these models depend on.

An accompanying policy memo, "AI News Audit: AI, Canadian Journalism, and Paths for Policy Action", discusses the implications of these findings for Canadian copyright, the Online News Act, statutory licensing, and attribution standards.

A note on AI in this research. This project was an experiment in developing an AI-assisted research methodology. Two senior researchers, Aengus Bridgman and Taylor Owen, designed a pipeline in which AI tools were embedded at each stage of the process, from study design to data collection and response coding to statistical analysis and prose drafting to graphic design, to test what this methodology could produce in a compressed timeline. Claude (Anthropic) was the primary AI tool used throughout. All code and content was reviewed, tested, and verified by Bridgman and/or Owen. The methodology itself is part of what this project set out to test.

Pourquoi c'est important

Le journalisme canadien fait partie de l'infrastructure démocratique. Il est essentiel à la reddition de comptes, au débat public et à la participation civique éclairée. Le journalisme est aussi coûteux à produire et sous de fortes pressions financières depuis deux décennies. Les entreprises d'IA en extraient désormais la valeur à chaque étape : en ingérant les archives de presse comme données d'entraînement, en produisant du contenu dérivé sans nommer les sources, et en livrant des réponses aux consommateurs qui peuvent réduire le besoin et l'incitatif de consulter la source originale.

Ce n'est pas une préoccupation hypothétique. Nos données montrent que les modèles d'IA possèdent une connaissance détaillée de l'actualité canadienne, une connaissance qui, pour la majorité des sujets testés, provient fort probablement du journalisme canadien. Ils répondent à des questions sur la politique, les politiques publiques et la société canadiennes avec une précision qui désigne le journalisme canadien comme source principale. Ils incluent souvent des liens vers des sites de médias canadiens, mais nomment rarement les journalistes ou les organisations dont les reportages alimentent la réponse. Un consommateur qui ne clique pas sur chaque lien ne saura pas quel journalisme il vient de consulter. Le résultat est un système dans lequel les entreprises d'IA monétisent le savoir des journalistes tandis que les salles de rédaction qui l'ont produit reçoivent des liens mais pas de reconnaissance. Cette dynamique, si elle se maintient à grande échelle, risque d'accélérer le déclin même du journalisme dont ces modèles dépendent.

Une note de politique complémentaire, « AI News Audit: AI, Canadian Journalism, and Paths for Policy Action », examine les implications de ces résultats pour le droit d'auteur canadien, la Loi sur les nouvelles en ligne, les licences légales et les normes d'attribution.

Note sur l'utilisation de l'IA dans cette recherche. Ce projet constituait une expérimentation en méthodologie de recherche assistée par l'IA. Deux chercheurs principaux, Aengus Bridgman et Taylor Owen, ont conçu un processus dans lequel des outils d'IA ont été intégrés à chaque étape – de la conception de l'étude à la collecte de données, du codage des réponses à l'analyse statistique, de la rédaction à la conception graphique – afin de tester ce que cette méthodologie pouvait produire dans un délai réduit. Claude (Anthropic) a été le principal outil d'IA utilisé tout au long du projet. L'ensemble du code et du contenu a été révisé, testé et vérifié par Bridgman et/ou Owen. La méthodologie elle-même fait partie de ce que ce projet cherchait à évaluer.

Context

The relationship between technology companies and news organizations has been contentious for over a decade. Social media platforms extracted advertising revenue while distributing news content, prompting regulatory responses including Canada's Online News Act (C-18) and Australia's News Media Bargaining Code.¹ AI companies present a different challenge: they ingest, process, and deliver journalists' work directly to consumers, often without visible attribution. Major AI companies have signed licensing deals with U.S. and U.K. publishers, but no Canadian news organization has reached a comparable agreement – the major Canadian publishers are suing instead.² The value transfer is asymmetric: journalists produce the knowledge that AI models draw on, but AI companies may capture the traffic and revenue that journalism needs to survive.

AI models interact with Canadian journalism at three distinct stages, each involving a different form of value extraction. At the **ingestion** stage, AI companies scrape and train on Canadian news archives, embedding journalists' knowledge into model weights (the internal parameters that encode everything a model has learned). At the **production** stage, AI models synthesize that journalism into substitute products: answers that deliver much of the reporting's value directly to consumers, with links to news sites but rarely naming the originating outlet.³ At the **distribution** stage, AI models become the channel through which the public encounters news, but rarely send readers back to the source, diverting the traffic that sustains journalism.

Each stage maps to a distinct policy question. Copyright law governs ingestion, and Canada's fair dealing doctrine under the Copyright Act has never been applied to large-scale AI training.³ Derivative works and competition law govern production: the creation of substitute goods from copyrighted material.⁴ Bargaining codes and platform regulation govern distribution, but C-18 was designed for platforms that link to news, not models that synthesize and deliver journalism's value in a fundamentally different way.⁵ The empirical question must come before the policy question: are AI models actually doing this? At what scale? With what consequences for which news organizations?

Against this backdrop, we designed an empirical audit of four leading AI companies and their use of Canadian journalism. This brief traces the three stages of that value chain – ingestion, production, and distribution.

Questions

1. Have AI models absorbed Canadian journalism?

We test whether the back catalogue of Canadian news is already embedded in AI training data (the text a model was trained on before release) by asking models about real events from the past three years without web access, particularly stories about domestic politics, provincial affairs, and local events that receive little or no international coverage.

2. Do AI models produce journalism's value without its costs?

When models draw on Canadian reporting to answer consumer questions, do they deliver something that substitutes for the original article? We measure how often AI responses cover enough distinctive reporting to reduce a consumer's need to visit the source, and whether paywalls provide any meaningful protection.

3. Do consumers get sent back to the source?

AI models are becoming a primary channel through which the public encounters news. We test whether consumers have any pathway back to the original reporting, and which outlets get visibility when models do cite sources.

Figure 1 summarizes the research design. Track 1 tested for knowledge of events from headlines derived from 2,267 Canadian news stories (English and French) using economy and flagship models from four AI companies – ChatGPT (OpenAI), Claude (Anthropic), Gemini (Google), and Grok (xAI) – with web search disabled, so any knowledge the models show must come from their training data. We use each company's disclosed training cutoff (the date after which no new information was included in the model's training): stories before the cutoff could have been ingested during training, while stories after it serve as a control. Track 2 enables web search and asks models about 140 specific February 2026 articles from both free and paywalled Canadian outlets, measuring whether the resulting responses substitute for the original journalism and whether Canadian news outlets are mentioned and/or users are directed back to the source. We coded each response using a combination of rule-based checks and LLM-based assessment (see Methodology for full details).

Methodology

Two complementary empirical tracks testing how AI agents interact with Canadian journalism

AGENTS TESTED



ChatGPT



Gemini

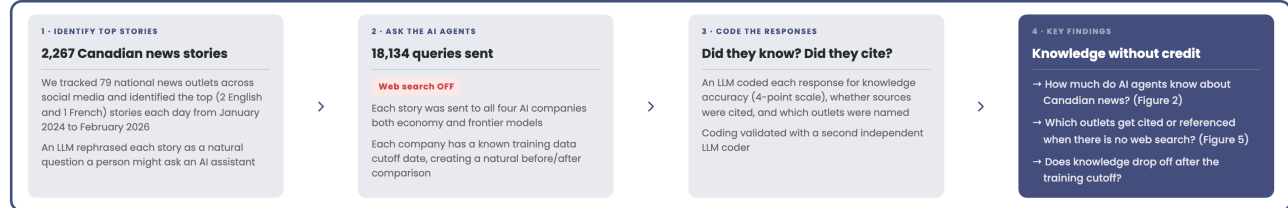


Claude

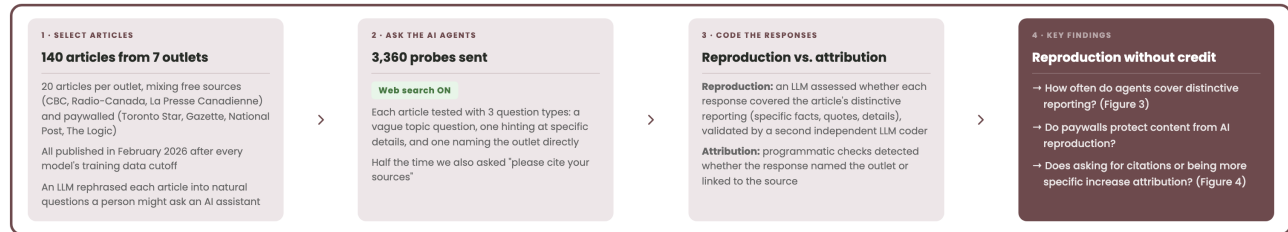


Grok

TRACK 1 Content Audit — Do AI agents know Canadian news stories? Do they credit the source?



TRACK 2 Attribution Audit — Do AI agents reproduce journalism without attribution? Even from behind paywalls?



Data collection: February 27–28, 2026 - Four AI companies tested: OpenAI (ChatGPT), Google (Gemini), Anthropic (Claude), xAI (Grok)
Full methodological details including model specifications, coding protocols, and intercoder reliability statistics are in the Methodology appendix.

Figure 1: Methodology overview: two empirical tracks testing how AI models interact with Canadian journalism.

Analysis

1) Have AI models absorbed Canadian journalism?

The first question is foundational: have AI companies ingested Canadian journalism into their training data? To test this, we prompted economy and flagship models from all four AI companies on 2,267 real Canadian news stories (English and French) between January 2024 and February 2026 **without enabling web search**. Disabling web access means any knowledge the models show must

come from their training data. Each model received a brief headline-style description of a real story and was asked what it knew. We then coded whether the response showed accurate knowledge of the event and whether it attributed that knowledge to any source. Stories before a model's training cutoff could have been ingested during training; stories after it could not. If a model knows pre-cutoff stories but not post-cutoff ones, the knowledge almost certainly comes from training data – and for these overwhelmingly domestic stories (Canadian politics, provincial policy, local events), Canadian journalism is the most plausible primary source.

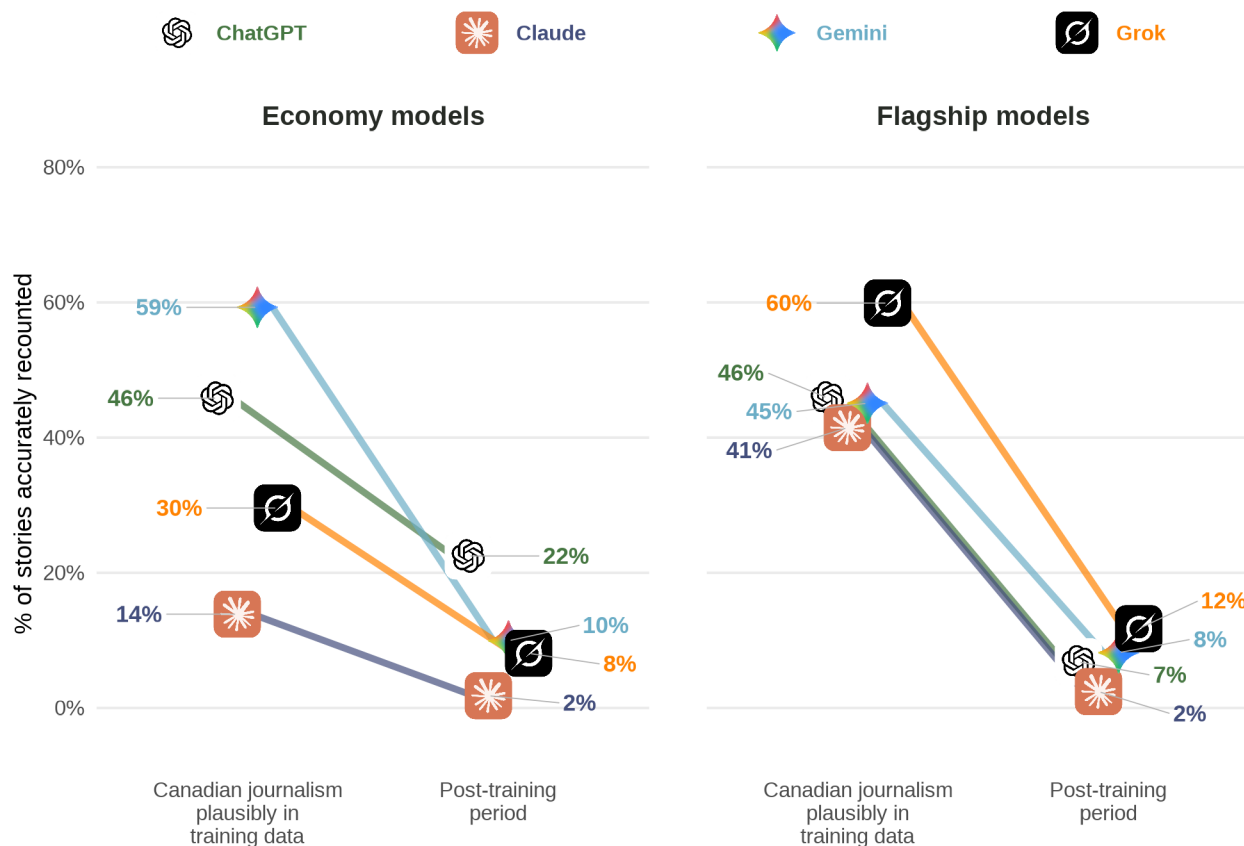


Figure 2: Accuracy of AI model responses to Canadian news stories, by training data cutoff. Economy and flagship models queried without web search. N = 18,134 responses across 2,267 stories (English and French).

Figure 2 shows the results. Note that our corpus draws exclusively from national outlets; local and regional coverage was not tested. The pattern is consistent across model tiers. Models showed at least partial knowledge of pre-cutoff Canadian news events in 76% of economy-tier responses and 76% of flagship-tier responses, then dropped sharply for post-cutoff stories (11% and 7% accurate respectively). Among knowledgeable responses, economy models were accurate 37% of the time on average; flagship models 48%. Flagship models score higher overall (unsurprisingly, given they typically have larger

training datasets) but the decline after the training cut-off appears across all models and tiers. This points to training data as the source of the knowledge rather than general reasoning, and suggests that more capable models have absorbed more Canadian news content, not less. The specific sources (Canadian journalism, wire services, Wikipedia, or other publicly available text covering the same events) cannot be isolated from model behavior alone.¹ ChatGPT and Grok show higher post-cutoff accuracy than Claude or Gemini, but this reflects hallucination – generating confident-sounding but fabricated in-

¹Economy-tier knowledge rates vary substantially by model: ChatGPT (89%), Gemini (94%), and Grok (90%) all demonstrated knowledge in the large majority of pre-cutoff responses, while Claude demonstrated knowledge in only 32% – not because it lacked training data, but because it defaults to refusing rather than guessing when uncertain. As a result, Claude's low rate pulls down the economy-tier average, causing it to understate the knowledge the other three models actually demonstrate. Two-sample t-tests comparing per-response accuracy between economy and flagship models: pre-training stories $t = 11.61$, $p < 0.001$; post-training stories $t = 7.34$, $p < 0.001$. Flagship models were 11.9 percentage points more accurate on pre-training stories and 4.6 percentage points less accurate on post-training stories. Note that each model has a different training data cutoff date, so the pre- and post-cutoff story pools differ in size and date range across models and tiers; the averages reported here are means of per-model percentages.

formation – not knowledge. Inspection of the responses reveals that these models confidently generate plausible-sounding but fabricated accounts of events they cannot have seen in training data: inventing specific but wrong outcomes (e.g., stating the Bank of Canada *held* rates when it actually *cut* them), producing generic descriptions vague enough to be unfalsifiable, or fabricating events wholesale. Claude, by contrast, typically acknowledges ignorance, resulting in lower measured “accuracy” but more honest responses.²

Despite showing substantial knowledge of Canadian current events, models rarely attributed it. Models showed at least partial knowledge in 74% of responses to pre-cutoff stories. Among those responses, 92% provided no source attribution of any kind – no named outlet, no recommendation to check a specific news organization, no general reference to “Canadian media.” The training cutoff dates used here are self-reported by each company and may be approximate; some information could also reach models through reinforcement learning from human feedback (RLHF) – a process where human evaluators help fine-tune model behavior after initial training – system prompt updates, or other post-training processes. The sharp accuracy drop visible in Figure 2 at each model’s reported cutoff is consistent with these dates being broadly accurate.

2) Do AI models produce journalism’s value without its costs?

AI models have absorbed Canadian journalism at scale and almost never attribute it, even for stories within their training window. We tested directly what happens when models can actively search the web. We enabled web

search and asked models about 140 real articles from February 2026. Figure 3 shows the default consumer experience: what happens when someone asks a generic topic question without requesting citations. This is how most people use AI models: “Tell me about X,” not “What did the Toronto Star report about X?”

With web search enabled, models retrieve content from across the web, including Canadian news sites, and synthesize it into responses. The blue squares show how often the result covers enough of the article’s distinctive reporting (specific events, named individuals, key findings) that a reader could plausibly get the gist of the story without visiting the news site. These are not complete reproductions: they are partial summaries and paraphrases that cover some of the original article’s distinctive content, though they sometimes contain factual errors or omissions (see “A note on accuracy” below). We evaluated each response against the source article to determine whether it covered the article’s distinctive reporting, not merely the general topic.³ The green squares show how often the model credits the source by naming the outlet in the response text or via structured machine-readable citations returned alongside the response.

Coverage rates are high while attribution rates are not. Gemini and Claude covered distinctive reporting in 81% and 72% of responses respectively, but Gemini credited the source only 6% of the time. Grok covered distinctive reporting in 59% of responses while citing the source in only 7% of them. ChatGPT, one of the most widely used models, covered distinctive content in 54% of responses but almost never credited the originating newsroom. Even when models fail to cover the distinctive reporting, they still deliver a topical response that can reduce the consumer’s motivation to visit the source. AI models use web search to retrieve current journalism and

²For post-cutoff stories, 72% of ChatGPT’s economy-tier responses addressed the topic substantively enough to be coded as showing at least partial knowledge – despite the model never having seen the stories in training. Of those, 69% were coded as inaccurate: the model generated confident, detailed responses about the right topic but fabricated the specifics (e.g., when asked about Calgary’s June 2024 water crisis caused by a catastrophic water main break, it invented a scenario involving flash flooding and overwhelmed storm infrastructure). Grok’s economy responses were similar: 68% addressed post-cutoff stories substantively, with 88% of those coded inaccurate. Claude generated substantive responses for only 7% of post-cutoff stories, declining to answer the rest. Flagship models are substantially less prone to this behaviour: ChatGPT’s flagship responses were coded as showing knowledge of only 24% of post-cutoff stories (vs. 72% economy) and Grok’s flagship 41% (vs. 68%), suggesting that hallucination is concentrated in cheaper, smaller models. This means the post-cutoff “accuracy” figures in Figure 2 likely overstate true knowledge for economy-tier ChatGPT and Grok, which hallucinate more freely, and understate the relative honesty of models like Claude that refuse to guess.

³The “covers distinctive reporting” measure combines three LLM-coded levels: partial coverage (56.8% of all generic-unprompted responses), close paraphrase (7.9%), and verbatim reproduction (1.6%). Partial coverage – where the response covers article-specific events and findings in paraphrased form rather than reproducing exact facts or language – accounts for the majority of viable substitutes. Rule-based pattern-matching alone captures only a fraction of this semantic coverage, which is why LLM-based assessment was necessary. Even partial coverage of a story’s distinctive reporting can reduce the consumer’s need to visit the source, but readers should note that the blue category reflects “covers some of the article-specific reporting,” not full reproduction.

deliver it directly to consumers, rarely crediting the news-

room that produced it.⁴

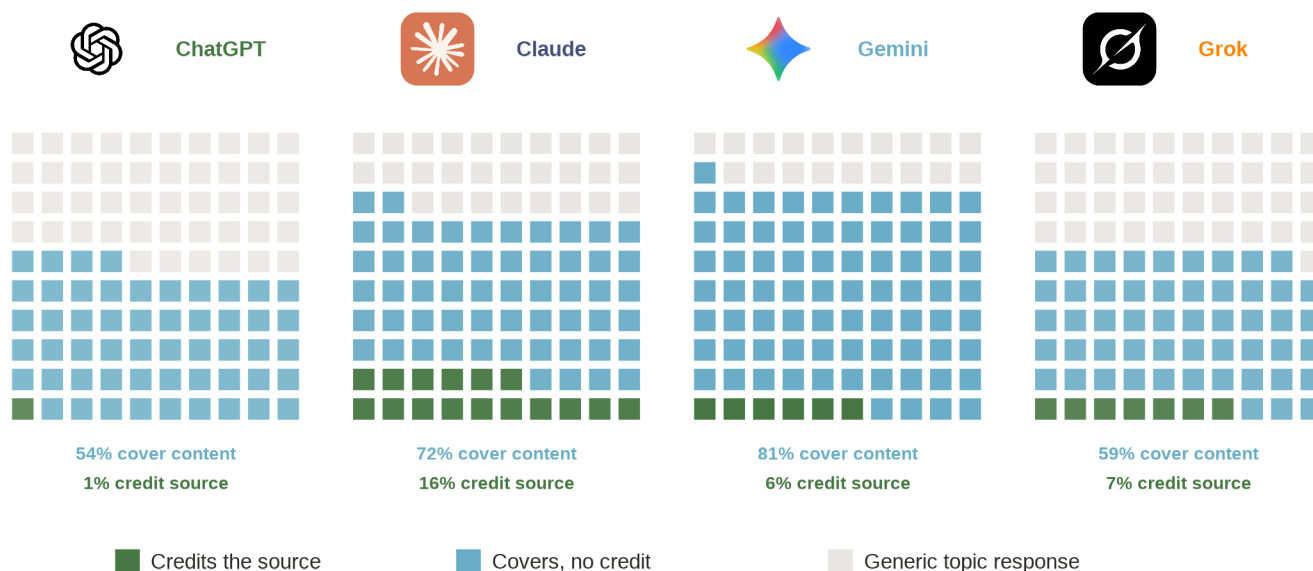


Figure 3: Coverage of distinctive reporting and source attribution for recent Canadian news articles. Economy models with web search enabled under generic framing with no citation prompt. N = 560 responses across 140 articles.

This substitution extends to paywalled content. We tested whether paywalls provide any protection across seven outlets spanning the spectrum: free (CBC News, Radio-Canada, Canadian Press), general paywalled (Toronto Star, Montreal Gazette, National Post), and niche paywalled (The Logic). Averaged across all paywalled outlets, models covered distinctive reporting from originally paywalled content (64%) at rates comparable to free outlets (70%). The competitive news market ensures the same information appears elsewhere – stories are syndicated through wire services, reposted on aggregators, and shared across social media. Paywalls hold only when an outlet’s journalism is distinctive enough that no free or more easily accessible substitute exists; competitive coverage renders most paywalls irrelevant.

For most stories, paywalls are bypassed rather than breached: the competitive news market ensures the same

information appears across multiple outlets, syndicated through wire services, reposted on aggregators, and shared on social media. However, API citation logs from some models showed direct citations to paywalled URLs, including hard-paywalled niche outlets, with extensive verbatim reproduction of the source text. This suggests that some models’ web retrieval tools may access paywalled content that would be blocked for human readers. Whether the mechanism is direct access or competitive syndication, the result is the same: paywalled journalism is reproduced without compensation.

A note on accuracy. Track 1 shows a clear pattern of hallucination for post-cutoff stories: models generate plausible-sounding narratives that match the general topic but get key facts wrong. When asked about the Bank of Canada’s June 2024 rate cut, ChatGPT’s economy model confidently reported that the Bank *held*

⁴A pilot sample of 407 flagship-tier responses showed no significant differences on any tested outcome (all $p > 0.05$; see Methodology). This is expected: substitution and attribution behaviour are driven by each provider’s search and citation architecture.

rates steady – the opposite of what happened. These responses are not retrieving stored knowledge; they are pattern-matching against generic templates (a funeral draws “thousands of mourners,” a cabinet shuffle follows “scandals and controversies”) and inventing specifics. Track 2 presents a more ambiguous picture. Among responses that covered distinctive reporting from the source article, 51% were coded as diverging from the original article’s reporting. This measure captures divergence from one specific article, not factual error – a response that accurately synthesizes several other sources covering the same story would also be flagged here. The substitution and attribution findings do not depend on this measure.

3) Do consumers get sent back to the source?

The third stage of the value chain is distribution: when AI models answer consumer questions, do those consumers have any pathway back to the journalism that informed the response? For publishers, this is the key question: do models credit and link to original source material?

Under realistic conditions (a generic topic question with no special request for citations), most models link to Canadian news sites at moderate-to-high rates but almost never name the originating outlet. Gemini (69%), Claude (66%), and Grok (45%) all provided Canadian news URLs in many cases, while ChatGPT did so less frequently (29%). But the more visible form of attribution (naming the outlet) was rare across all models, ranging from 1% to 16% – a pattern consistent with independent audits of AI search citation practices.⁶ Models often include a link to Canadian journalism, but without naming the source in the response text, consumers who do not click through every URL will not know whose reporting informed the answer. Beyond Canadian news URLs, all models link at high rates to non-Canadian sources such as Wikipedia, Reuters, AP, or U.S. news outlets, none of which generate traffic for the Canadian newsrooms whose journalism created the content.

Figure 4 tests whether AI models can be pushed toward better attribution. We varied both the framing of the question and whether consumers explicitly asked models to cite sources. The columns show three framing conditions: 1) generic topic questions (as in Figure 3); 2) questions referencing specific facts from the article; and 3) questions

that directly name the outlet. The rows show whether the prompt requested citations. Each panel plots the four models on two dimensions: the probability of naming the outlet (x-axis) and the probability of linking to a Canadian news site (y-axis). The upper-right quadrant (green shading) is where models both name and link – the closest equivalent to social media posts sharing news content with clear attribution and a click-through URL. The lower-left (orange shading) is the worst case: models neither name the outlet nor link to a Canadian news site, leaving consumers with no path back to the source and news organizations with no way to recapture value from their work.

Under the most favourable conditions (directly naming the outlet and explicitly asking for citations), attribution improves substantially across all models. All four named the outlet in a majority of responses: Claude (97%), Gemini (95%), ChatGPT (86%), and Grok (74%). Linking rates were also strong: Grok (91%), Gemini (69%), Claude (64%), and ChatGPT (59%). Meaningful attribution is technically achievable. The gap between the default experience and the best-case scenario is a core finding: most consumers will never explicitly name an outlet or ask for citations, so the generic-condition results reflect the experience that shapes the market for journalism.

The distribution consequences extend beyond individual articles. Outlet mentions across all 18,134 Track 1 responses fall into two categories: outlets named as the source of a specific claim, and outlets recommended when a model cannot answer. Both patterns reveal the same concentration dynamic, but they reflect different model behaviours.

Among outlets used as genuine sources, CBC News leads with 238 attributions, followed by The Globe and Mail (93), Radio-Canada (89), La Presse (52), and Global News (44) – driven almost entirely by Grok, which accounts for 95% of all named-as-source mentions. CTV News, despite its size, received only 26 named-as-source mentions. The Toronto Star received 10; the Montreal Gazette, 1.

The recommended category is far larger, with 6,845 mentions total, but reflects a different phenomenon: models pointing users elsewhere when they lack knowledge. CBC News (1,972), The Globe and Mail (1,206), and CTV News (1,149) dominate here, driven overwhelmingly by Claude, which accounts for 71% of all recommended mentions. These are deflections, however, and cluster

around the same small set of nationally prominent free outlets. The Toronto Star received 133 recommended

mentions; Financial Post, 25; Montreal Gazette, 3.

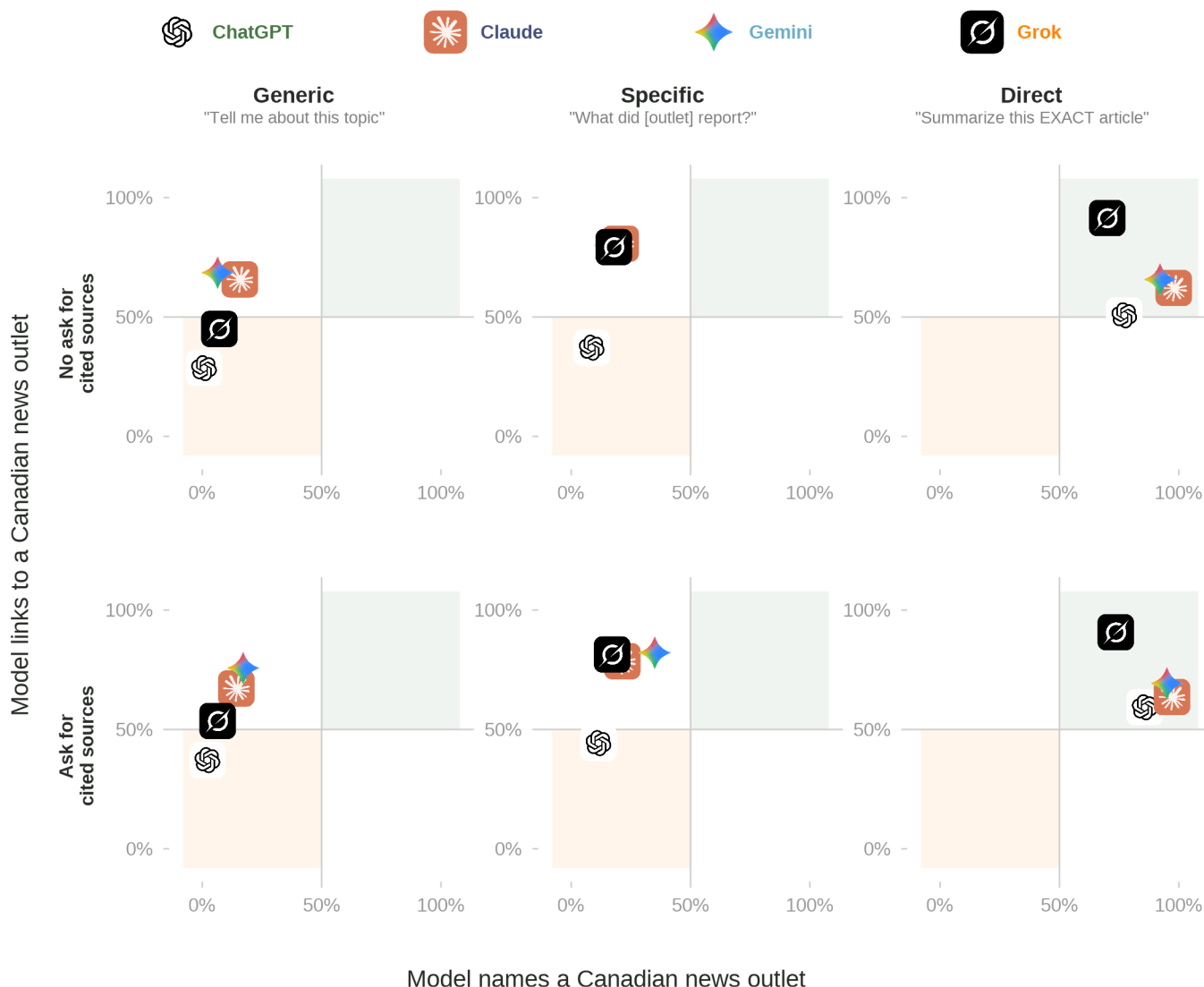


Figure 4: Source naming and linking rates by framing condition and citation prompt. Economy models with web search enabled. N = 3,360 responses across 140 articles.

AI models surface the outlets consumers already know, while smaller and paywalled outlets that may produce significant original reporting are largely absent. Among English-language outlets, CBC, CTV, and Global News – all freely accessible – capture the most AI visibility in both categories. The Globe and Mail performs relatively well,

but the Toronto Star and Financial Post are marginal despite being important newsrooms. Regional Postmedia papers serving Calgary, Edmonton, Ottawa, and Vancouver are essentially absent. Among French-language outlets, Radio-Canada and La Presse dominate, with Le Devoir a distant third. The Journal de Montréal, one of Que-

bec's most widely read papers, received only 48 total mentions across all models. Figure 5 plots total AI mentions against monthly web visits for each outlet, combining outlets named as sources ($n = 761$) and outlets recommended for further reading ($n = 6,845$). The diagonal reference line shows proportional representation where

an outlet's AI mentions match its share of total readership. Canadian outlets tend to sit above this line: when answering questions about Canadian news, AI models cite Canadian sources more than their global audience share would predict, which is the expected pattern.

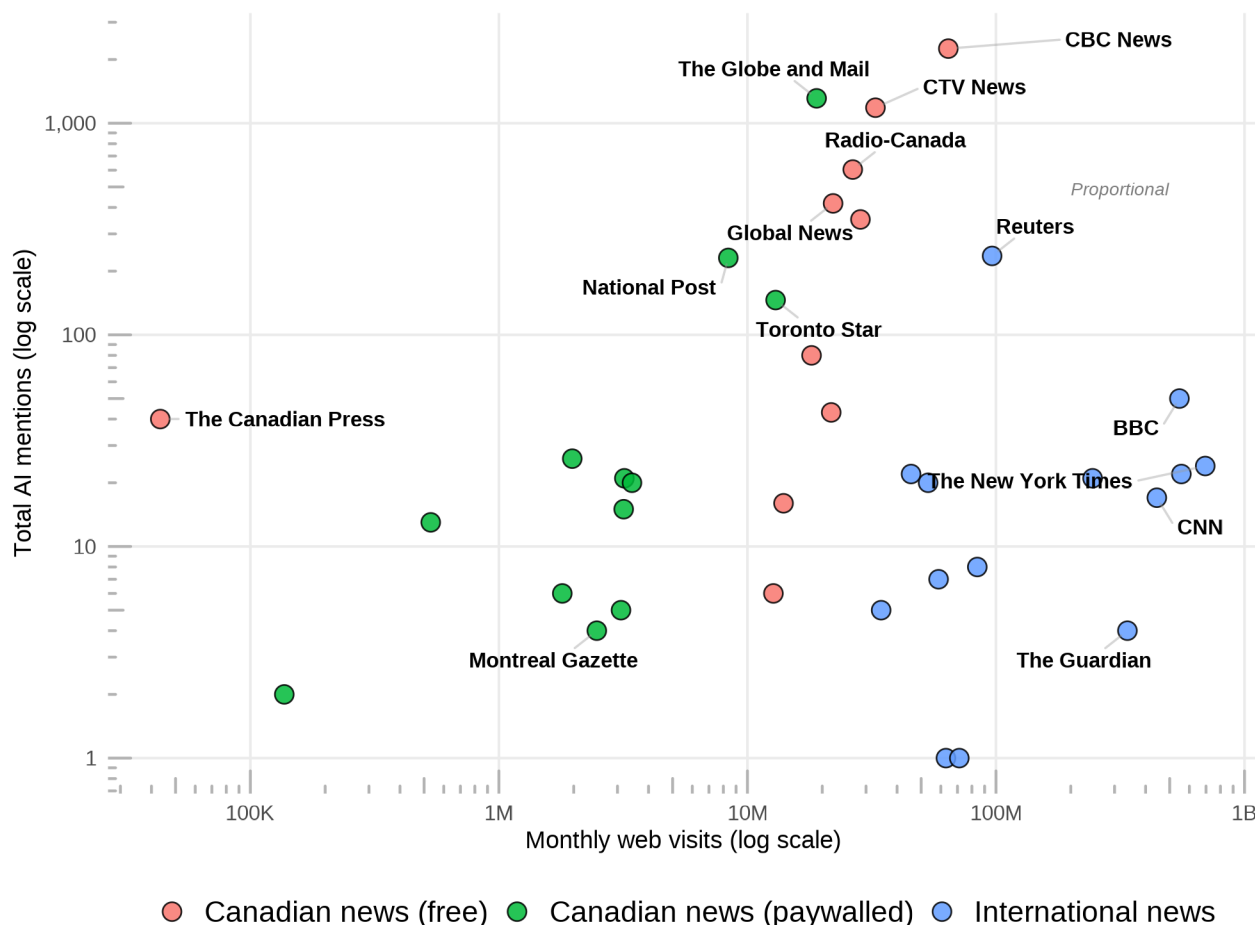


Figure 5: Total AI mentions vs. monthly web visits by outlet. Mentions include outlets named as sources ($n = 761$) and outlets recommended for further reading ($n = 6,845$). Economy and flagship models queried without web search. $N = 18,134$ responses (English and French). Readership data from SimilarWeb (January 2026).

The distribution of citations was statistically different between economy and flagship model tiers (economy mean = 0.59, flagship mean = 0.21 Canadian sources cited per response; $p < 0.001$). AI models also cited 807

non-news sources a total of 2,249 times, focusing predominantly on Canadian government agencies (Statistics Canada, Global Affairs Canada, Bank of Canada), polling firms (Nanos, Angus Reid, Leger), and institutional

⁵Models differ substantially in citation style. In the economy tier (English and French combined), Grok named a source inline in 20% of responses, the highest of any model, and accounts for 95% of all named-as-source outlet mentions across all tiers. Claude almost never names

sources. These are omitted from Figure 5 to focus on news outlets.⁵

Across all 9,227 responses to pre-cutoff stories (English and French combined), 82% cited no source whatsoever. The pattern holds across both language groups: 81% of

responses to French-language stories provided no citation, compared to 82% for English-language stories. The attribution gap is not a language-specific phenomenon.⁶ French-language journalism is doubly disadvantaged: its content is absorbed into model training data, but the outlets that produced it are almost never acknowledged.

Conclusion

AI models possess extensive knowledge of Canadian current events consistent with having ingested Canadian news reporting (though the exact provenance cannot be isolated from model behavior alone, as noted above). They answer detailed questions about Canadian politics, economics, and society with a specificity most consistent with having trained on Canadian journalism. With web access, they go further, producing responses that cover enough of the original reporting that could reduce a consumer's need to visit the source (often with factual errors the consumer has no easy way to catch, because source names are rarely provided in the response text). The competitive structure of news coverage means most Canadian journalism is freely discoverable somewhere, and AI models capture that value while rarely identifying the source in the response itself.

Under default API conditions (approximating unprompted consumer use), most models included links to Canadian news sites in a majority of responses. But they rarely named any Canadian source in the response text itself: a consumer who does not click through every

link will not know whose journalism informed the answer. When consumers explicitly asked for citations, all models improved substantially. Naming sources is technically achievable today; the failure is one of default design, not capability.⁸ The result is a system in which AI companies extract the informational value of journalism – the facts, analysis, and context that journalists produce – while the newsrooms that created it receive links but not recognition in the text consumers actually read.

The path forward requires no new technical capability: meaningful attribution is a choice. Canada has set precedent through the Online News Act: technology companies profiting from journalistic content must share value with producers. AI companies represent a more direct form of value extraction than social media platforms: they absorb journalism, transform it, and redistribute it without identifying its origin. But improved citation alone will not resolve the underlying tension. The current news model depends on human journalists interacting with democratic actors, and AI-mediated distribution severs the economic link between the production and consumption of journalism. Regulatory frameworks may need to go beyond attribution to address the value transfer from Canadian journalism to AI companies.⁷

sources inline (under 1% of responses); instead it recommends outlets at the end of a response when it cannot answer, accounting for 71% of all recommended mentions and typically listing the same short roster of nationally prominent free outlets (CBC, Globe and Mail, CTV). ChatGPT almost exclusively recommends rather than names sources (98% of its outlet mentions are recommendations), but does so infrequently: only 8% of its economy-tier responses mention any outlet at all. Gemini rarely attributes sources regardless of knowledge level (97% of economy responses cite nothing), answering from training data without acknowledging where information came from. These behavioural differences mean cross-model citation counts are not directly comparable.

⁶When French-language outlets were cited, they appeared in only 10% of responses to French stories. The disparity holds across all four models and both tiers.

Methodology

This study uses two complementary empirical tracks. Track 1 (Content Audit) tests whether AI models have absorbed Canadian journalism into their training data and whether they attribute it. We prompted four major AI models on 2,267 real Canadian news stories with web search disabled, so any knowledge they demonstrate must come from training data. Track 2 (Attribution Audit) tests what happens when models can actively search the web: we enabled web search and asked the same four models about 140 specific articles from seven Canadian outlets, using a factorial design that varies how directly the question references the source. Together, the two tracks cover the full pipeline from ingestion to production to distribution. Figure 1 provides a visual overview; the subsections below detail each step.

Story corpus construction

We constructed a corpus of the top Canadian news stories by mining social media posts from 79 national Canadian news outlets tracked by the Media Ecosystem Observatory across Facebook, X/Twitter, Instagram, YouTube, TikTok, and Bluesky. For each day between January 1, 2024 and February 15, 2026 (777 days), we:

1. Extracted posts from national news outlets, separating English-language and French-language posts;
2. Filtered to posts containing Canadian-specific keywords (provinces, cities, prominent political figures including Trudeau, Poilievre, Carney, Ford, and institutions including the Bank of Canada, RCMP, and Parliament);
3. Applied TF-IDF vectorization and agglomerative clustering (cosine distance, threshold = 0.7) separately to each language stream to identify distinct stories;
4. Ranked stories within each language by outlet breadth (number of distinct outlets covering the story) and total engagement;
5. Selected the top two English-language stories and the top one French-language story per day.

This produced 1,511 distinct English-language Canadian news stories and 756 French-language stories (2,267 total), spanning federal politics, provincial affairs, economic policy, social issues, and more. Each story record includes the representative headline, the set of outlets that covered it, total engagement (likes, shares, comments), and

engagement as a share of all news engagement that day.

AI models tested

We tested the four major AI models consumers interact with daily using each company's economy-tier model (the cheaper, faster version designed for high-volume use) via API (application programming interface – the programmatic access point that lets researchers query models systematically) to ensure standardized, reproducible queries at scale: ChatGPT (OpenAI, GPT-5-mini, cutoff May 2024), Grok (xAI, Grok 4 Fast, cutoff November 2024), Gemini (Google, Gemini 3 Flash, cutoff January 2025), and Claude (Anthropic, Claude Haiku 4.5, cutoff February 2025). We queried all models on February 27–28, 2026. API access replicates underlying model behavior but not every feature of consumer-facing interfaces (such as custom system prompts or built-in search features); see Consumer product validation below. For Track 1, we also tested each company's flagship model (see Robustness check below).

The knowledge cutoff dates are critical for interpreting Track 1 results: stories occurring after a model's cutoff cannot have been absorbed during training. These dates are self-reported by each company and may be approximate; some information could also leak through reinforcement learning from human feedback (RLHF), system prompt updates, or other post-training processes. However, the sharp accuracy drop observed in Figure 2 at each model's reported cutoff is consistent with the dates being broadly accurate. Our study period (January 2024 – February 2026) spans a range where some stories fall within every model's training window, some fall within only the more recent models' windows, and the most recent stories fall outside all models' cutoffs, creating a natural experiment in how training data recency affects knowledge and attribution.

Track 1: Content audit design

Objective: Test the extent to which AI models possess knowledge of Canadian current events and whether they spontaneously attribute Canadian news organizations.

Prompt construction: For each of the 2,267 stories, we used a local language model (Qwen3-30B) to rephrase the story headlines into natural, open-ended questions in the story's language (English or French) – the kind a person might casually ask an AI assistant. All questions be-

gin with “What,” “Why,” “How,” or “Who” (never yes/no) and include a time reference (e.g., “in March 2024”) to anchor the query. Well-known public figures and places are preserved for specificity, but outlet names were stripped from headlines before rephrasing to avoid priming the AI models toward any particular news organization.

Example prompts:

- “What happened with Air Canada’s on-time performance in January 2024 that made it rank last in North America?”
- “What led Rachel Notley to step down as Alberta’s NDP leader in January 2024 after nearly ten years at the helm?”
- “What happened in January 2024 when Canada announced a two-year cap on international student admissions?”
- «Qu’est-ce qui s’est passé à Montréal en janvier 2024 lorsqu’un délit de fuite a coûté la vie à deux piétons dans Ahuntsic-Cartierville?»

Execution: We sent each prompt to all eight models without web search enabled. This is the critical design choice: by disabling web access, we ensure that any knowledge the models demonstrate about these events must come from their training data. Because the stories tested are overwhelmingly domestic in nature, Canadian journalism is the most plausible source for this knowledge. The system prompt was minimal and language-matched: “You are a helpful assistant. Answer the user’s question about Canadian news. Be concise.” (English) and “Tu es un assistant utile. Réponds à la question de l’utilisateur sur l’actualité canadienne. Sois concis.” (French).

Total queries: 2,267 stories (1,511 English + 756 French) × 8 models (4 economy + 4 flagship) = 18,134 responses.

Track 2: Attribution audit design

Objective: Test whether AI models produce viable substitutes for Canadian journalism, including content originally published behind paywalls, when consumers ask about current events.

Article corpus: We obtained articles published in February 2026 from Nexis Uni across seven Canadian outlets spanning the spectrum from free to hard-paywalled, and from national English-language to French-language and niche outlets:

- Free: CBC News, Radio-Canada, La Presse Canadienne (Canadian Press)
- Paywalled (general): Toronto Star, Montreal Gazette, National Post
- Paywalled (niche): The Logic (technology and business policy)

The selection of publisher was based on availability of data from Nexis Uni; future audits should cover a wider range of stories collected directly from the outlets themselves. All articles are from February 2026 and post-date every model’s training data cutoff, ensuring that any content substitution requires active web retrieval and not training data recall.

Article selection: We selected 20 Canada-focused articles per outlet (140 total), prioritized by length (>800 words), original reporting depth, and topic diversity. For each article, we used Claude Haiku 4.5 to identify five to eight distinctive facts: named individuals (especially non-public figures like interviewees and witnesses), specific dollar amounts, direct quotes, or unique details that function as fingerprints of the source article and to generate the three framing-level prompts. Five initial articles were hand-crafted to validate the automated pipeline; the remainder were generated programmatically. The inclusion of The Logic as a niche paywalled outlet provides variation in paywall type (hard vs. soft) and content distinctiveness.

Factorial design (3 × 2 × 4 × N):

Factor 1 Framing (three levels):

- *Generic (F1)*: A natural consumer question about the topic, with no reference to the article or outlet (e.g., “What’s happening with suburban malls and condo development in the Greater Toronto Area?”).
- *Specific (F2)*: A question referencing distinctive facts from the article, without naming the outlet (e.g., “I heard a condo project at a Toronto-area mall was cancelled because almost no units sold...something like less than 10% presold”).
- *Direct (F3)*: A question explicitly naming the outlet or headline (e.g., “Can you tell me about the Toronto Star’s recent reporting on suburban mall condo cancellations?”).

Factor 2 Citation prompt (two levels):

- *Unprompted (C0)*: No instruction about sources.

- *Prompted (C1):* Appends “Please cite your sources.”

Factor 3 Model:

Economy-tier models from all four AI companies, with *web search enabled*, reflected the default consumer experience. Each provider’s native search capability was activated: OpenAI’s web search preview, Google’s Search grounding, Anthropic’s web search tool, and xAI’s web search.

Factor 4 Articles:

Total queries: The full factorial design calls for $3 \times 2 \times 4 \times 140 = 3,360$ probes. Economy models completed all 3,360 probes.

Execution: We ran each query in a fresh conversation (no context bleed between conditions) via each provider’s API on February 27–28, 2026. All four providers were queried in parallel (independent rate limits), with 1-second delays between sequential calls to the same provider. For each response, we recorded the full text, token usage (input/output), API-level citations where available, estimated cost, and elapsed time. Output tokens were uncapped (16,384 max) to avoid truncating model responses.

Coding and analysis

With 18,134 Track 1 responses and 3,360 Track 2 probes, manual coding was infeasible. We used a two-layer coding approach combining rule-based and LLM-based assessment.

For Track 2, rule-based coding provided the foundation for source citation and URL detection: rule-based source citation detection (checking whether the originating outlet is named in the response text or in structured API-level citation metadata) and Canadian news URL matching against 76 domains. We also applied rule-based fact fingerprinting (checking whether pre-identified distinctive facts appear in responses) and verbatim sequence matching (four-or-more-word sequence overlap between article text and model response). Claude and Gemini return structured grounding citations alongside their responses: URLs and titles of sources their search systems retrieved. We count a model as “citing the source” if the originating outlet appears in either the response text or these API-level citations; in approximately 5% of responses, the API metadata referenced the source outlet

even though the response text did not mention it. These rule-based measures require no subjective judgment and are fully reproducible.

For the central question of whether a response functions as a viable substitute for the original article, we relied on LLM-based coding using GPT-5.2 (OpenAI, via Batch API) and Qwen3.5-35B-A3B-FP8. For each Track 2 response, the coding model received three inputs: (1) the full source article – outlet, headline, date, paywall status, pre-identified distinctive facts, and the complete article text; (2) the probe prompt, including its framing and citation condition; and (3) the AI agent’s response. For Track 1, the inputs were simpler: the consumer prompt, the agent’s response, and ground truth metadata (original headline and which Canadian outlets covered the story). In both cases, the LLM assessed whether the response covered the article’s distinctive reporting or merely discussed the general topic. This distinction is critical: a response about “Canadian immigration policy” that discusses general trends is not a substitute for a specific Toronto Star investigation, while one that covers the article’s central findings is. Rule-based coding alone cannot capture this distinction because AI models routinely paraphrase rather than copy, producing viable substitutes without triggering verbatim or exact-fact matches. The substitution rates reported in Figure 3 use the LLM-coded assessment; Figure 4’s citation pathway analysis uses rule-based coding for source naming and URL detection. Using one company’s model (GPT-5.2) to evaluate all four companies’ outputs – including OpenAI’s own ChatGPT – is a potential source of bias; see Validation below.

Track 1 coding schema. We coded each response on four dimensions:

- Knowledge level (`knowledgeable` / `partial` / `no_knowledge` / `refusal`): Does the model demonstrate factual knowledge of the event? “Knowledgeable” requires specific, correct details; “partial” indicates hedging or vague accuracy; “no_knowledge” means the model clearly lacks information or is entirely wrong.
- Citation type (`named_as_source` / `recommended` / `vague_reference` / `none`): How does the model reference sources? We distinguish between using an outlet as an information source (“According to CBC...”), recommending outlets for the user to check (“I’d suggest looking at the Globe and Mail”), vague attribution (“Canadian media reported...”), and no attribution at

all.

- Sources cited: Every news outlet or media organization mentioned by name, classified as Canadian or non-Canadian.
- Accuracy (*accurate* / *mostly_accurate* / *inaccurate* / *unverifiable*): Factual correctness compared against the original headline and ground truth.

Track 2 coding schema. We coded each response on four dimensions: content reproduction level, distinctive fact count, attribution level, and link quality. Content reproduction level is an ordinal scale: *verbatim* (word-for-word reproduction of article passages), *close_paraphrase* (key content reproduced with minor rewording), *partial* (article-specific events and findings covered in paraphrased form, but not the full story), *topic_only* (general topic discussed without distinctive reporting), and *none* (no content overlap). We collapse the first three into “covers distinctive reporting” for main analyses. Under the generic unprompted condition, the breakdown was: *partial* 56.8%, *close paraphrase* 7.9%, *verbatim* 1.6%. The remaining dimensions: number of pre-identified distinctive facts reproduced (0–8 score); attribution level (full, outlet only, vague, none, misattribution); and link quality (working, broken, hallucinated, none).

Canadian news URL detection. Figure 4 distinguishes between models that include *any* URL in their responses and those that link specifically to a Canadian news site. We define “links to a Canadian news site” as the presence in either the response text or structured API-level citation metadata of a URL or domain reference matching one of 76 Canadian news domains. For Gemini, which returns citations as redirect URLs through Google’s grounding infrastructure with the destination domain in the citation title field, we match against the title metadata in addition to the URL itself. This list includes major national outlets (CBC, Globe and Mail, CTV, Global News, National Post, Toronto Star, La Presse), regional dailies (Winnipeg Free Press, Calgary Herald, Vancouver Sun, Ottawa Citizen, Times Colonist, SaltWire), Postmedia chain papers (Toronto Sun, Montreal Gazette, St. Albert Gazette), Indigenous news outlets (APTN, Nunatsiaq News), political and policy outlets (iPolitics, The Hub, Policy Options, Western Standard, The Walrus), Yahoo News Canada subdomains, public broadcasting (TVO, Radio-Canada), and Canadian specialty publications (Sportsnet, BetaKit, BlogTO, Toronto Life). The distinction matters: many models include URLs at high rates (e.g., ChatGPT in-

cludes a URL in over 90% of responses), but these links often point to non-Canadian sources such as Wikipedia, Reuters, AP, or U.S. news outlets, none of which generate traffic for the originating Canadian newsrooms whose journalism informed the AI’s response.

Validation. We validated the LLM coding pipeline using two approaches. First, we manually reviewed a random sample of 100 coded responses across all four models. Citation type classifications showed high agreement between human and LLM judgments; disagreements were concentrated at the boundary between “partial” and “knowledgeable” knowledge levels. Second, we independently re-coded all Track 1 (18,134 responses) and Track 2 (3,438 responses) using Qwen3.5-35B-A3B-FP8 (Alibaba Cloud), a separate large language model from a different model family, running on Digital Research Alliance infrastructure on an NVIDIA H100 GPU with identical prompts. We report quadratic-weighted kappa (κ_w) for ordinal dimensions, which credits partial agreement on adjacent categories. All kappas are computed on the pooled corpus (English + French, economy + flagship). The measures that drive our main findings – citation type ($\kappa_w = 0.85$) and Canadian source count ($\kappa_w = 0.94$) – achieved almost perfect agreement. The ordinal knowledge-level scale reached substantial weighted agreement ($\kappa_w = 0.70$); Track 2 reproduction level reached substantial agreement ($\kappa_w = 0.78$); and the binary paywalled-content measure achieved strong agreement ($\kappa_w = 0.80$). Factual accuracy was the weakest dimension in both tracks (Track 1 $\kappa_w = 0.49$; Track 2 $\kappa_w = 0.50$), consistent with its inherently subjective nature. Disagreements on the reproduction and knowledge scales were almost entirely adjacent-category errors rather than fundamental misclassifications.

Robustness check: flagship models

A potential concern with our design is the use of each company’s economy-tier model. To address this, we tested whether results hold when using a flagship model (the most capable and expensive version) from each company: GPT-5.2 (OpenAI, cutoff August 2025), Gemini 3 Pro (Google, cutoff January 2025), Grok 4 (xAI, cutoff November 2024), and Claude Sonnet 4.6 (Anthropic, cutoff March 2025).

Track 1 (full replication and inclusion). Track 1’s low per-query cost made it feasible to re-query all 2,267 stories

under identical conditions (same prompts, same system message, no web search). In some cases, flagship models have later knowledge cutoff dates than the economy-tier models (e.g., GPT-5.2 is trained to August 2025 vs. GPT-5-mini's May 2024).

Track 2 (pilot only). Track 2's web-search-enabled design is substantially more expensive per query, so we collected a pilot sample of 407 flagship-tier responses and compared them against the economy-tier results on 4 key outcomes (viable substitute rate, source naming, Canadian URL provision, paywall substitution). Of 4 two-sample t-tests, 1 reached significance at $p < 0.05$ uncorrected and 1 after Holm correction for 4 comparisons (reproduction: $p = 0.083$; source naming: $p = 0.082$; paywall substitution: $p = 0.65$), with Canadian URL rates also non-significant ($p = < 0.001$). This is expected: Track 2's dominant factor is each provider's search and citation architecture, which is shared across model tiers.

Robustness check: consumer product validation

A second concern is that consumer-facing products (e.g., chatgpt.com, claude.ai) may behave differently from API endpoints. To assess this, we manually tested five Track 1 prompts through the consumer web interfaces of ChatGPT and Claude. We selected the five stories for diversity across the timeline: two pre-cutoff (Pearson Airport gold heist, April 2024; Poilievre non-confidence motion, March 2024), one near-cutoff (K'naan allegations, September 2024), and two post-cutoff (Air Canada flight attendant strike, August 2025; Trump altered map, January 2026). For each, we compared the consumer response against our API flagship and economy responses for the same prompt.

Consumer products differed from API responses in response quality: consumer answers were typically three to five times longer and more detailed but the core citation finding held: across all ten consumer tests (five prompts $\times$ two platforms), zero named a specific Canadian news outlet as a source. In one notable case (Air Canada strike), the API flagship GPT-5.2 recommended five Canadian outlets by name (CBC, Globe and Mail, CTV, Global News, Canadian Press), but consumer ChatGPT's response on the identical prompt cited none. Consumer Claude recommended checking "CBC News, Globe and Mail" on that story, consistent with Claude's API-level re-

ferral behavior, but did not cite them as sources for its answer.

Consumer products also exhibited greater willingness to generate confident responses about post-cutoff events without web search. On the Trump map story (January 2026), both API flagships correctly declined to answer, while both consumer products generated detailed responses with consumer ChatGPT producing approximately 350 words of plausible but unverified content, and consumer Claude providing a more accurate account including political context not available in any API tier. These results suggest that the consumer experience, if anything, amplifies the patterns documented in our API-based analysis: richer derivative content with equal or less attribution to original sources.

We also tested five Track 2 prompts through the same consumer interfaces with web search enabled, selecting articles for diversity: two from French-language outlets (La Presse Canadienne, Radio-Canada), two from paywalled English outlets (Toronto Star, The Logic), and one from a free English outlet (CBC). For Track 2, consumer and API responses were comparable in quality: the search and citation architecture matters more than the model tier. In this small sample, neither the Toronto Star's nor The Logic's distinctive reporting was reproduced by consumer responses. When consumer Claude found facts from the Star's reporting (Cloverdale Mall's cancelled condo project), it attributed them to a free secondary source (Toronto Life) not the Star. The two French-language articles tested were poorly served: neither consumer product found the Radio-Canada article, and the La Presse Canadienne article was discovered only through English-language syndication. Consumer ChatGPT with web search did cite Canadian outlets effectively (e.g., The Tyee, Global News, CityNews Montreal, MBC Radio on an Indigenous rights story), and consumer Claude provided the strongest Canadian-source citation behavior of any configuration tested.

Consumer web interfaces cannot be queried programmatically at scale, making a systematic replication of all 2,267 prompts impractical. This spot-check is necessarily limited, but the consistency of the patterns across ten stories, two platforms, and both tracks provides reasonable confidence that our API-based findings are consistent with consumer-facing products.

Cost and reproducibility

The total API cost for the entire study, including data collection across all four AI providers and LLM-based response coding, was approximately \$252 USD. Of this, OpenAI accounted for \$115 (data collection + all LLM coding via Batch API), Anthropic \$59, xAI \$40, and Google \$38. This cost covers 6,044 Track 1 economy queries, 6,044 Track 1 flagship queries, 3,360 Track 2 economy probes, 407 Track 2 flagship probes, and the automated coding of all responses. We executed all queries on February 27–28, 2026. Prompts, response data, and analysis code are available from the authors on request.

Limitations

This study was designed iteratively rather than pre-registered; analytical choices including the coding schema and reproduction threshold were refined during the analysis phase. Several additional design limitations should be noted beyond those outlined in the Context section above. First, our paywall findings are ambiguous: API citation logs from some models included direct references to paywalled URLs with extensive verbatim reproduction, but we cannot definitively determine whether models accessed the content through the paywall, through cached or indexed versions, or through other freely available sources. Second, our Track 2 article corpus is modest in size (140 articles across seven outlets), which limits our ability to draw outlet-level conclu-

sions; the paywall analysis in particular compares 240 free-outlet probes against 320 paywalled-outlet probes. Third, Track 2 uses economy-tier models only; a pilot sample of 407 flagship-tier responses showed no significant differences on key outcomes (see Robustness check). Fourth, our manual validation sample for both tracks ($N = 100$) is small relative to the dataset; while the primary findings rest on rule-based measures, the LLM-coded dimensions (knowledge level, accuracy) would benefit from more extensive human validation. Fifth, our primary data collection uses API endpoints rather than consumer-facing interfaces; a manual spot-check found consumer products exhibit equal or worse citation behavior (see Consumer product validation), but the consumer experience may still differ. Sixth, we did not test Google Search's AI Overview, which is arguably among the most common ways consumers now encounter AI-generated summaries of news; AI Overviews use the same underlying Gemini model but likely with different prompting and retrieval architecture, and typically include source links and hoverable attribution that may produce different citation behavior than the standalone chatbot interface we tested. Seventh, our Track 1 story corpus selects the top stories each day by outlet breadth and engagement, biasing the sample toward the most prominent national events; results may not generalize to lower-salience or local journalism where training data availability will be more limited. Finally, AI models are updated frequently; our results reflect a point-in-time audit (February 27–28, 2026) and may not generalize to future versions.

Endnotes

1. Canada's Online News Act (C-18) received Royal Assent in June 2023. Australia's News Media Bargaining Code was enacted in 2021. In October 2024, the CRTC approved Google's exemption application under C-18, committing \$100 million annually to Canadian news organizations – the first major implementation of the Act. <https://www.canada.ca/en/canadian-heritage/services/online-news.html>; CRTC, "CRTC Approves Google's Application and Paves Way for Annual \$100 Million Contribution to Canadian News Organizations," October 2024, <https://www.canada.ca/en/radio-television-telecommunications/news/2024/10/crtc-approves-googles-application-and-paves-way-for-annual-100-million-contribution-to-canadian-news-organizations.html>
2. OpenAI, Google, and Meta have signed licensing deals with major U.S. and U.K. publishers including the Washington Post, the Guardian, the Associated Press, and News Corp, with some deals reportedly worth up to \$50 million per year. No Canadian news organization has reached a comparable agreement. Instead, the major Canadian publishers – Toronto Star Newspapers Limited, the Globe and Mail, Postmedia, Canadian Press, and CBC/Radio-Canada – filed suit against OpenAI in Ontario Superior Court in November 2024 (Case No. CV-24-00732231-00CL). In November 2025, the court rejected OpenAI's jurisdictional challenge, allowing the case to proceed in Canada. For a timeline of deals, see Digiday, "A Timeline of the Major Deals Between Publishers and AI Tech Companies in 2025," 2025, <https://digiday.com/media/a-timeline-of-the-major-deals-between-publishers-and-ai-tech-companies-in-2025/>; for a comprehensive tracker, see the Tow Center for Digital Journalism, "AI Deals and Lawsuits," Columbia Journalism Review, <https://tow.cjr.org/ai-deals-lawsuits/>; for the statement of claim, see <https://litigate.com/assets/uploads/Canadian-News-Media-Companies-v-OpenAI.pdf>
3. The Government of Canada held a consultation on copyright in the age of generative AI in 2023 but produced no legislative action. The Artificial Intelligence and Data Act (AIDA), originally Part 3 of Bill C-27, died when Parliament was prorogued in January 2025. As of March 2026, Canada has no AI-specific legislation in force. See ISED Canada, "Consultation on Copyright in the Age of Generative Artificial Intelligence: What We Heard Report," 2024, <https://ised-isde.canada.ca/site/strategic-policy-sector/en/marketplace-framework-policy/consultation-copyright-age-generative-artificial-intelligence-what-we-heard-report>; <https://www.parl.ca/legisinfo/en/bill/44-1/c-27>
4. In *Advance Local Media LLC v. Cohere Inc.* (N.D.N.Y., November 2025), a U.S. federal court ruled that AI-generated "substitutive summaries" of news articles may plausibly infringe copyright, even when non-verbatim, because they reduce the market for the original work. The ruling has not been tested in Canada, but it signals growing judicial recognition that AI-generated derivative content can constitute a substitute for the original journalism.
5. The Online News Act defines "digital news intermediaries" as platforms that make news content available through an online service, a definition built around linking and display, not synthesis and summarization. <https://www.parl.ca/DocumentViewer/en/44-1/bill/1/C-18/royal-assent>
6. A March 2025 audit by Columbia Journalism Review's Tow Center compared eight AI search engines and found that all failed to accurately cite news sources, with chatbots frequently directing users to syndicated copies rather than original reporting. See Jaźwińska, K. and Chandrasekar, A., "AI Search Has a Citation Problem," Columbia Journalism Review, March 6, 2025. https://www.cjr.org/tow_center/we-compared-eight-ai-search-engines-theyre-all-bad-at-citing-news.php
7. The Artificial Intelligence and Data Act (AIDA), originally Part 3 of Bill C-27, died when Parliament was prorogued in January 2025. The government has signaled it will address AI copyright issues separately but has introduced no replacement legislation as of March 2026. <https://www.parl.ca/legisinfo/en/bill/44-1/c-27>
8. The tendency of large language models to reproduce information from training data without citing original sources is a well-documented limitation of current architectures. See Chen et al., "Benchmarking Large Language Models in Retrieval-Augmented Generation," *Proceedings of the AAAI Conference on Artificial Intelligence*, 2024. <https://arxiv.org/abs/2307.02185>

Contributors

Aengus Bridgman is Associate Director, Research, and Director of the Media Ecosystem Observatory at the Centre for Media, Technology and Democracy, McGill University.

Taylor Owen is the Beaverbrook Chair in Media, Ethics and Communication and Founding Director of the Centre for Media, Technology and Democracy, McGill University.

This brief was made possible by the considerable work of a large number of researchers in Canada. We would especially like to thank Alexei Abrahams, Esli Chan, Mika Desblancs-Patel, David Hobson, Diya Jiang, Mathieu Lavigne, Saewon Park, Zeynep Pehlivan, Christopher Ross, Ben Steels, and Ashley Vu for building and maintaining the Media Ecosystem Observatory digital trace data collection pipeline.

The Media Ecosystem Observatory is a research initiative based at McGill University that monitors and analyzes Canada's digital information ecosystem. The Observatory is part of the Centre for Media, Technology and Democracy, which empowers democratic societies to navigate digital change. This research was conducted through the Canadian Digital Media Research Network (CDMRN), a national research network focusing on digital media challenges in Canada. The work of the Observatory and the CDMRN is funded in part by Heritage Canada's Digital Citizen Initiative.

The work of the Media Ecosystem Observatory and the Canadian Digital Media Research Network is funded in part by Heritage Canada's Digital Citizen Initiative.

Canada 