

Sixty years ago the futurist Arthur C. Clarke observed that any sufficiently advanced technology is indistinguishable from magic. The internet—how we both communicate with one another and together preserve the intellectual products of human civilization—fits Clarke’s observation well. In Steve Jobs’s words, “it just works,” as readily as clicking, tapping, or speaking. And every bit as much aligned with the vicissitudes of magic, when the internet doesn’t work, the reasons are typically so arcane that explanations for it are about as useful as trying to pick apart a failed spell. Underpinning our vast and simple-seeming digital networks are technologies that, if they hadn’t already been invented, probably wouldn’t unfold the same way again. They are artifacts of a very particular circumstance, and it’s unlikely that in an alternate timeline they would have been designed the same way.

The internet’s distinct architecture arose from a distinct constraint and a distinct freedom: First, its academically minded designers didn’t have or expect to raise massive amounts of capital to build the network; and second, they didn’t want or expect to make money from their invention. The internet’s framers thus had no money to simply roll out a uniform centralized network the way that, for example, FedEx metabolized a capital outlay of tens of millions of dollars to deploy liveried planes, trucks, people, and drop-off boxes, creating a single point-to-point delivery system. Instead, they settled on the equivalent of rules for how to bolt existing networks together.

Rather than a single centralized network modeled after the legacy telephone system, operated by a government or a few massive utilities, the internet was designed to allow any device anywhere to interoperate with any other device, allowing any provider able to bring whatever networking capacity it had to the growing party. And because the network’s creators did not mean to monetize, much less monopolize, any of it, the key was for desirable content to be provided naturally by the network’s users, some of whom would act as content producers or hosts, setting up watering holes for others to frequent.

Unlike the briefly ascendant proprietary networks such as CompuServe, AOL, and Prodigy, content and network would be separated. Indeed, the internet had and has no main menu, no CEO, no public stock offering, no formal organization at all. There are only engineers who meet every so often to refine its suggested communications protocols that hardware and software makers, and network builders, are then free to take up as they please.

So the internet was a recipe for mortar, with an invitation for anyone, and everyone, to bring their own bricks. Tim Berners-Lee took up the invite and invented the protocols for the World Wide Web, an application to run on the internet. If your computer spoke “web” by running a browser, then it could speak with servers that also spoke web, naturally enough known as websites. Pages on sites could contain links to all sorts of things that would, by definition, be but a click away, and might in practice be found at servers anywhere else in the world, hosted by people or organizations not only not affiliated with the linking webpage, but entirely unaware of its existence. And webpages themselves might be assembled from multiple sources before they displayed as a single unit, facilitating the rise of ad networks that could be called on by websites to insert surveillance beacons and ads on the fly, as pages were pulled together at the moment someone sought to view them.

And like the internet’s own designers, Berners-Lee gave away his protocols to the world for free—enabling a design that omitted any form of centralized management or control, since there was no usage to track by a World Wide Web, Inc., for the purposes of billing. The web, like the internet, is a collective hallucination, a set of independent efforts united by common technological protocols to appear as a seamless, magical whole. This absence of central control, or even easy central monitoring, has long been celebrated as an instrument of grassroots democracy and freedom. It’s not trivial to censor a network as organic and decentralized as the internet. But more recently, these features have been understood to facilitate vectors for individual harassment and societal destabilization, with no easy gating points through which to remove or label malicious work not under the umbrellas of the major social-media platforms, or to quickly identify their sources. While both assessments have power to them, they each gloss over a key feature of the distributed web and internet: Their designs naturally create gaps of responsibility for maintaining valuable content that others rely on. Links work seamlessly until they don’t. And as tangible counterparts to online work fade, these gaps represent actual holes in humanity’s knowledge.

Before today’s internet, the primary way to preserve something for the ages was to consign it to writing—first on stone, then parchment, then papyrus, then 20-pound acid-free paper, then a tape drive, floppy disk, or hard-drive platter—and store the result in a temple or library: a building designed to guard it against rot, theft, war, and natural disaster. This approach has facilitated preservation of some material for thousands of years. Ideally, there would be multiple identical copies stored in multiple libraries, so the failure of one storehouse wouldn’t extinguish the knowledge within. And in rare instances in which a document was surreptitiously altered, it could be compared against copies elsewhere to detect and correct the change. These buildings didn’t run themselves, and they weren’t mere warehouses. They were staffed with clergy and then librarians, who fostered a culture of preservation and its many elaborate practices, so precious documents would be both safeguarded and made accessible at scale—certainly physically, and, as important, through careful indexing, so an inquiring mind could be paired

with whatever a library had that might slake that thirst. (As Jorge Luis Borges pointed out, a library without an index becomes paradoxically less informative as it grows.)

At the dawn of the internet age, 25 years ago, it seemed the internet would make for immense improvements to, and perhaps some relief from, these stewards' long work. The quirkiness of the internet and web's design was the apotheosis of ensuring that the perfect would not be the enemy of the good. Instead of a careful system of designation of "important" knowledge distinct from day-to-day mush, and importation of that knowledge into the institutions and cultures of permanent preservation and access (libraries), there was just the infinitely variegated web, with canonical reference websites like those for academic papers and newspaper articles juxtaposed with PDFs, blogs, and social-media posts hosted here and there. Enterprising students designed web crawlers to automatically follow and record every single link they could find, and then follow every link at the end of that link, and then build a concordance that would allow people to search across a seamless whole, creating search engines returning the top 10 hits for a word or phrase among, today, more than 100 trillion possible pages. As Google puts it, "The web is like an ever-growing library with billions of books and no central filing system."

Now, I just quoted from Google's corporate website, and I used a hyperlink so you can see my source. Sourcing is the glue that holds humanity's knowledge together. It's what allows you to learn more about what's only briefly mentioned in an article like this one, and for others to double-check the facts as I represent them to be. The link I used points to <https://www.google.com/search/howsearchworks/crawling-indexing/>. Suppose Google were to change what's on that page, or reorganize its website anytime between when I'm writing this article and when you're reading it, eliminating it entirely. Changing what's there would be an example of content drift; eliminating it entirely is known as link rot.

It turns out that link rot and content drift are endemic to the web, which is both unsurprising and shockingly risky for a library that has "billions of books and no central filing system." Imagine if libraries didn't exist and there was only a "sharing economy" for physical books: People could register what books they happened to have at home, and then others who wanted them could visit and peruse them. It's no surprise that such a system could fall out of date, with books no longer where they were advertised to be—especially if someone reported a book being in someone else's home in 2015, and then an interested reader saw that 2015 report in 2021 and tried to visit the original home mentioned as holding it. That's what we have right now on the web.

Whether humble home or massive government edifice, hosts of content can and do fail. For example, President Barack Obama signed the Affordable Care Act in the spring of 2010. In the fall of 2013, congressional Republicans shut down day-to-day government funding in an attempt to kill Obamacare. Federal agencies, obliged to cease all but essential activities, pulled the plug on websites across the U.S. government, including access to thousands, perhaps millions, of official government documents, both current and archived, and of course very few having anything to do with Obamacare. As night follows day, every single link pointing to the affected documents and sites no longer worked.

In 2010, Justice Samuel Alito wrote a concurring opinion in a case before the Supreme Court, and his opinion linked to a website as part of the explanation of his reasoning. Shortly after the opinion was released, anyone following the link wouldn't see whatever it was Alito had in mind when writing the opinion. Instead, they would find this message: "Aren't you glad you didn't cite to this webpage ... If you had, like Justice Alito did, the original content would have long since disappeared and someone else might have come along and purchased the domain in order to make a comment about the transience of linked information in the internet age." Inspired by cases like these, some colleagues and I joined those investigating the extent of link rot in 2014 and again this past spring.

The first study, with Kendra Albert and Larry Lessig, focused on documents meant to endure indefinitely: links within scholarly papers, as found in the *Harvard Law Review*, and judicial opinions of the Supreme Court. We found that 50 percent of the links embedded in Court opinions since 1996, when the first hyperlink was used, no longer worked. And 75 percent of the links in the *Harvard Law Review* no longer worked. People tend to overlook the decay of the modern web, when in fact these numbers are extraordinary—they represent a comprehensive breakdown in the chain of custody for facts. Libraries exist, and they still have books in them, but they aren't stewarding a huge percentage of the information that people are linking to, including within formal, legal documents. No one is. The flexibility of the web—the very feature that makes it work, that had it eclipse CompuServe and other centrally organized networks—diffuses responsibility for this core societal function.

The problem isn't just for academic articles and judicial opinions. With John Bowers and Clare Stanton, and the kind cooperation of *The New York Times*, I was able to analyze approximately 2 million externally facing links found in articles at nytimes.com since its inception in 1996. We found that 25 percent of deep links have rotted. (Deep links are links to specific content—think theatlantic.com/article, as opposed to just theatlantic.com.) The older the article, the less likely it is that the links work. If you go back to 1998, 72 percent of the links are dead. Overall, more

than half of all articles in *The New York Times* that contain deep links have at least one rotted link.

Our studies are in line with others. As far back as 2001, a team at Princeton University studied the persistence of web references in scientific articles, finding that the raw number of URLs contained in academic articles was increasing but that many of the links were broken, including 53 percent of those in the articles they had collected from 1994. Thirteen years later, six researchers created a data set of more than 3.5 million scholarly articles about science, technology, and medicine, and determined that one in five no longer points to its originally intended source. In 2016, an analysis with the same data set found that 75 percent of all references had drifted.

Of course, there's a keenly related problem of permanency for much of what's online. People communicate in ways that feel ephemeral and let their guard down commensurately, only to find that a Facebook comment can stick around forever. The upshot is the worst of both worlds: Some information sticks around when it shouldn't, while other information vanishes when it should remain. So far, the rise of the web has led to routinely cited sources of information that aren't part of more formal systems; blog entries or casually placed working papers at some particular web address have no counterparts in the pre-internet era. But surely anything truly worth keeping for the ages would still be published as a book or an article in a scholarly journal, making it accessible to today's libraries, and preservable in the same way as before? Alas, no.

Because information is so readily placed online, the incentives for creating paper counterparts, and storing them in the traditional ways, declined slowly at first and have since plummeted. Paper copies were once considered originals, with any digital complement being seen as a bonus. But now, both publisher and consumer—and libraries that act in the long term on behalf of their consumer patrons—see digital as the primary vehicle for access, and paper copies are deprecated. From my vantage point as a law professor, I've seen the last people ready to turn out the lights at the end of the party: the law-student editors of academic law journals. One of the more stultifying rites of passage for entering law students is to “subcite,” checking the citations within scholarship in progress to make sure they are in the exacting and byzantine form required by legal-citation standards, and, more directly, to make sure the source itself exists and says what the citing author says it says. (In a somewhat alarming number of instances, it does not, which is a good reason to entertain the subciting exercise.)

The original practice for, say, the *Harvard Law Review*, was to require a student subciter to lay eyes on an original paper copy of the cited source, such as a statute or a judicial opinion. The Harvard Law Library would, in turn, endeavor to keep a physical copy of everything—ideally every law and case from everywhere—for just that purpose. The *Law Review* has since eased up, allowing digital images of printed text to suffice, and that’s not entirely unwelcome: It turns out that the physical law (as distinct from the laws of physics) takes up a lot of space, and Harvard Law School was sending more and more books out to a remote depository, to be laboriously retrieved when needed.

A few years ago I helped lead an effort to digitize all of that paper both as images and as searchable text—more than 40,000 volumes comprising more than 40 million pages—which completed the scanning of nearly every published case from every state from the time of that state’s inception up through the end of 2018. (The scanned books have been sent to an abandoned limestone mine in Kentucky, as a hedge against some kind of digital or even physical apocalypse.) A special quirk allowed us to do that scanning, and to then treat the longevity of the result as seriously as we do that of any printed material: American case law is not copyrighted, because it’s the product of judges. (Indeed, any work by the U.S. government is required by statute to be in the public domain.) But the Harvard Law School library is no longer collecting the print editions from which to scan—it’s too expensive. And other printed materials are essentially trapped on paper until copyright law is refined to better account for digital circumstances.

Into that gap has entered material that’s born digital, offered by the same publishers that would previously have been selling on printed matter. But there’s a catch: These officially sanctioned digital manifestations of material have an asterisk next to their permanence. Whether it’s an individual or a library acquiring them, the purchaser is typically buying mere access to the material for a certain period of time, without the ability to transfer the work into the purchaser’s own chosen container. This is true of many commercially published scholarly journals, for which “subscription” no longer signifies a regular delivery of paper volumes that, if canceled, simply means no more are forthcoming. Instead, subscription is for ongoing access to the entire corpus of journals hosted by the publishers themselves. If the subscription arrangement is severed, the entire oeuvre becomes inaccessible.

Libraries in these scenarios are no longer custodians for the ages of anything, whether tangible or intangible, but rather poolers of funding to pay for fleeting access to knowledge elsewhere.

Similarly, books are now often purchased on Kindles, which are the Hotel Californias of digital devices: They enter but can't be extracted, except by Amazon. Purchased books can be involuntarily zapped by Amazon, which has been known to do so, refunding the original purchase price. For example, 10 years ago, a third-party bookseller offered a well-known book in Kindle format on Amazon for 99 cents a copy, mistakenly thinking it was no longer under copyright. Once the error was noted, Amazon—in something of a panic—reached into every Kindle that had downloaded the book and deleted it. The book was, fittingly enough, George Orwell's *1984*. (*You don't have 1984. In fact, you never had 1984. There is no such book as 1984.*)

At the time, the incident was seen as evocative but not truly worrisome; after all, plenty of physical copies of *1984* were available. Today, as both individual and library book buying shifts from physical to digital, a de-platforming of a Kindle book—including a retroactive one—can carry much more weight. Deletion isn't the only issue. Not only can information be removed, but it also can be changed. Before the advent of the internet, it would have been futile to try to change the contents of a book after it had been long published. Librarians do not take kindly to someone attempting to rip out or mark up a few pages of an “incorrect” book. The closest approximation of post-hoc editing would have been to influence the contents of a later edition.

Ebooks don't have those limitations, both because of how readily new editions can be created and how simple it is to push “updates” to existing editions after the fact. Consider the experience of Philip Howard, who sat down to read a printed edition of *War and Peace* in 2010. Halfway through reading the brick-size tome, he purchased a 99-cent electronic edition for his Nook e-reader: As I was reading, I came across this sentence: “It was as if a light had been Nookd in a carved and painted lantern ...” Thinking this was simply a glitch in the software, I ignored the intrusive word and continued reading. Some pages later I encountered the rogue word again. With my third encounter I decided to retrieve my hard cover book and find the original (well, the translated) text. For the sentence above I discovered this genuine translation: “It was as if a light had been kindled in a carved and painted lantern ...”

A search of this Nook version of the book confirmed it: Every instance of the word *kindle* had been replaced by *nook*, in perhaps an attempt to alter a previously made Kindle version of the book for Nook use. Here are some screenshots I took at the time: It is only a matter of time before the retroactive malleability of these forms of publishing becomes a new area of pressure and regulation for content censorship. If a book contains a passage that someone believes to be defamatory, the aggrieved person can sue over it—and receive monetary damages if they're right. Rarely is the book's existence itself called into question, if only because of the difficulty of

putting the cat back into the bag after publishing. Now it's far easier to make demands for a refinement or an outright change of the offending sentence or paragraph. So long as those remedies are no longer fanciful, the terms of a settlement can include them, as well as a promise not to advertise that a change has even been made. And a lawsuit need never be filed; only a demand made, publicly or privately, and not one grounded in a legal claim, but simply one of outrage and potential publicity. Rereading an old Kindle favorite might then become reading a slightly (if momentarily) tweaked version of that old book, with only a nagging feeling that it isn't quite how one remembers it.

This isn't hypothetical. This month, the best-selling author Elin Hilderbrand published a new novel. The novel, widely praised by critics, included a snippet of dialogue in which one character makes a wry joke to another about spending the summer in an attic on Nantucket, "like Anne Frank." Some readers took to social media to criticize this moment between characters as anti-Semitic. The author sought to explain the character's use of the analogy before offering an apology and saying that she had asked her publisher to remove the passage from digital versions of the book immediately.

There are sufficient technical and typographical alterations to ebooks after they're published that a publisher itself might not even have a simple accounting of how often it, or one of its authors, has been importuned to alter what has already been published. Nearly 25 years ago I helped Wendy Seltzer start a site, now called Lumen, that tracks requests for elisions from institutions ranging from the University of California to the Internet Archive to Wikipedia, Twitter, and Google—often for claimed copyright infringements found by clicking through links published there. Lumen thus makes it possible to learn more about what's missing or changed from, say, a Google web search, because of outside demands or requirements.

For example, thanks to the site's record-keeping both of deletions and of the source and text of demands for removals, the law professor Eugene Volokh was able to identify a number of removal requests made with fraudulent documentation—nearly 200 out of 700 "court orders" submitted to Google that he reviewed turned out to have been apparently Photoshopped from whole cloth. The Texas attorney general has since sued a company for routinely submitting these falsified court orders to Google for the purpose of forcing content removals. Google's relationship with Lumen is purely voluntary—YouTube, which, like Google, has the parent company Alphabet, is not currently sending notices. Removals through other companies—like book publishers and distributors such as Amazon—are not publicly available.

The rise of the Kindle points out that even the concept of a link—a “uniform resource locator,” or URL—is under great stress. Since Kindle books don’t live on the World Wide Web, there’s no URL pointing to a particular page or passage of them. The same goes for content within any number of mobile apps, leaving people to trade screenshots—or, as *The Atlantic*’s Kaitlyn Tiffany put it, “the gremlins of the internet”—as a way of conveying content.

Here, courtesy of the law professor Alexandra Roberts, is how a district-court opinion pointed to a TikTok video: “A May 2020 TikTok video featuring the Reversible Octopus Plushies now has over 1.1 million likes and 7.8 million views. The video can be found at Girlfriends mood #teeturtle #octopus #cute #verycute #animalcrossing #cutie #girlfriend #mood #inamood #timeofmonth #chocolate #fyp #xyzcba #cbzzyz #t (tiktok.com).” Which brings us full circle to the fact that long-term writing, including official documents, might often need to point to short-term, noncanonical sources to establish what they mean to say—and the means of doing that is disintegrating before our eyes (or worse, entirely unnoticed). And even long-term, canonical sources such as books and scholarly journals are in fugacious configurations—usually to support digital subscription models that require scarcity—that preclude ready long-term linking, even as their physical counterparts evaporate.

The project of preserving and building on our intellectual track, including all its meanderings and false starts, is thus falling victim to the catastrophic success of the digital revolution that should have bolstered it. Tools that could have made humanity’s knowledge production available to all instead have, for completely understandable reasons, militated toward an ever-changing “now,” where there’s no easy way to cite many sources for posterity, and those that are citable are all too mutable. Again, the stunning success of the improbable, eccentric architecture of our internet came about because of a wise decision to favor the good over the perfect and the general over the specific. I have admiringly called this the “Procrastination Principle,” wherein an elegant network design would not be unduly complicated by attempts to solve every possible problem that one could imagine materializing in the future. We see the principle at work in Wikipedia, where the initial pitch for it would seem preposterous: “We can generate a consummately thorough and mostly reliable encyclopedia by allowing anyone in the world to create a new page and anyone else in the world to drop by and revise it.”

It would be natural to immediately ask what would possibly motivate anyone to contribute constructively to such a thing, and what defenses there might be against edits made ignorantly or in bad faith. If Wikipedia garnered enough activity and usage, wouldn't some two-bit vendor be motivated to turn every article into a spammy ad for a Rolex watch? Indeed, Wikipedia suffers from vandalism, and over time, its sustaining community has developed tools and practices for dealing with it that didn't exist when Wikipedia was created. If they'd been implemented too soon, the extra hurdles to starting and editing pages might have deterred many of the contributions that got Wikipedia going to begin with. The Procrastination Principle paid off.

Similarly, it wasn't on the web inventor Tim Berners-Lee's mind to vet proposed new websites according to any standard of truth, reliability, or ... anything else. People could build and offer whatever they wanted, so long as they had the hardware and connectivity to set up a web server, and others would be free to visit that site or ignore it as they wished. That websites would come and go, and that individual pages might be rearranged, was a feature, not a bug. Just as the internet could have been structured as a big CompuServe, centrally mediated, but wasn't, the web could have had any number of features to better assure permanence and sourcing. Ted Nelson's Xanadu project contemplated all that and more, including "two-way links" that would alert a site every time someone out there chose to link to it. But Xanadu never took off.

As procrastinators know, later doesn't mean never, and the benefits of the internet and web's flexibility—including permitting the building of walled app gardens on top of them that reject the idea of a URL entirely—now come at great risk and cost to the larger tectonic enterprise to, in Google's early words, "organize the world's information and make it universally accessible and useful." Sergey Brin and Larry Page's idea was a noble one—so noble that for it to be entrusted to a single company, rather than society's long-honed institutions, such as libraries, would not do it justice. Indeed, when Google's founders first released a paper describing the search engine they had invented, they included an appendix about "advertising and mixed motives," concluding that "the issue of advertising causes enough mixed incentives that it is crucial to have a competitive search engine that is transparent and in the academic realm." No such transparent, academic competitive search engine exists in 2021. By making the storage and organization of information everyone's responsibility and no one's, the internet and web could grow, unprecedentedly expanding access, while making any and all of it fragile rather than robust in many instances in which we depend on it.

What are we going to do about the crisis we're in? No one is more keenly aware of the problem of the internet's ephemerality than Brewster Kahle, a technologist who founded the Internet Archive in 1996 as a nonprofit effort to preserve humanity's knowledge, especially and including

the web. Brewster had developed a precursor to the web called WAIS, and then a web-traffic-measurement platform called Alexa, eventually bought by Amazon. That sale put Brewster in a position personally to help fund the Internet Archive's initial operations, including the Wayback Machine, specifically designed to collect, save, and make available webpages even after they've gone away. It did this by picking multiple entry points to start "scraping" pages—saving their contents rather than merely displaying them in a browser for a moment—and then following as many successive links as possible on those pages, and those pages' linked pages.

It is no coincidence that a single civic-minded citizen like Brewster was the one to step up, instead of our existing institutions. In part that's due to potential legal risks that tend to slow down or deter well-established organizations. The copyright implications of crawling, storing, and displaying the web were at first unsettled, typically leaving such actions either to parties who could be low key about it, saving what they scraped only for themselves; to large and powerful commercial parties like search engines whose business imperatives made showing only the most recent, active pages central to how they work; or to tech-oriented individuals with a start-up mentality and little to lose. An example of the latter is at work with Clearview AI, where a single rakish entrepreneur scraped billions of images and tags from social-networking sites such as Facebook, LinkedIn, and Instagram in order to build a facial-recognition database capable of identifying nearly any photo or video clip of someone.

Brewster is superficially in that category, too, but—in the spirit of the internet and web's inventors—is doing what he's doing because he believes in his work's virtue, not its financial potential. The Wayback Machine's approach is to save as much as possible as often as possible, and in practice that means a lot of things every so often. That's vital work, and it should be supported much more, whether with government subsidy or more foundation support. (The Internet Archive was a semifinalist for the MacArthur Foundation's "100 and Change" initiative, which awards \$100 million individually to worthy causes.)

A complementary approach to "save everything" through independent scraping is for whoever is creating a link to make sure that a copy is saved at the time the link is made. Researchers at the Berkman Klein Center for Internet & Society, which I co-founded, designed such a system with an open-source package called Amberlink. The internet and the web invite any form of additional building on them, since no one formally approves new additions. Amberlink can run on some web servers to make it so that what's at the end of a link can be captured when a webpage on an Amberlink-empowered server first includes that link. Then, when someone clicks on a link on an Amber-tuned site, there's an opportunity to see what the site had captured at that

link, should the original destination no longer be available. (Search engines such as Google have this feature, too—you can often ask to see the search engine’s “cached” copy of a webpage linked from a search-results page, rather than just following the link to try to see the site yourself.)

Amber is an example of one website archiving another, unrelated website to which it links. It’s also possible for websites to archive themselves for longevity. In 2020, the Internet Archive announced a partnership with a company called Cloudflare, which is used by popular or controversial websites to be more resilient against denial-of-service attacks conducted by bad actors that could make the sites unavailable to everyone. Websites that enable an “always online” service will see their content automatically archived by the Wayback Machine, and if the original host becomes unavailable to Cloudflare, the Internet Archive’s saved copy of the page will be made available instead.

These approaches work generally, but they don’t always work specifically. When a judicial opinion, scholarly article, or editorial column points to a site or page, the author tends to have something very distinct in mind. If that page is changing—and there’s no way to know if it will change—then a 2021 citation to a page isn’t reliable for the ages if the nearest copy of that page available is one archived in 2017 or 2024. Taking inspiration from Brewster’s work, and indeed partnering with the Internet Archive, I worked with researchers at Harvard’s Library Innovation Lab to start Perma. Perma is an alliance of more than 150 libraries. Authors of enduring documents—including scholarly papers, newspaper articles, and judicial opinions—can ask Perma to convert the links included within them into permanent ones archived at <http://perma.cc>; participating libraries treat snapshots of what’s found at those links as accessions to their collections, and undertake to preserve them indefinitely.

In turn, the researchers Martin Klein, Shawn Jones, Herbert Van de Sompel, and Michael Nelson have honed a service called Robustify to allow archives of links from whatever source, including Perma, to be incorporated into new “dual-purpose” links so that they can point to a page that works in the moment, while also offering an archived alternative if the original page fails. That could allow for a rolling directory of snapshots of links from a variety of archives—a networked history that is both prudently distributed, internet-style, while shepherded by the long-standing institutions that have existed for this vital public-interest purpose: libraries.

A technical infrastructure through which authors and publishers can preserve the links they draw on is a necessary start. But the problem of digital malleability extends beyond the technical. The law should hesitate before allowing the scope of remedies for claimed infringements of rights—whether economic ones such as copyright or more personal, dignitary ones such as defamation—to expand naturally as the ease of changing what’s already been published increases.

Compensation for harm, or the addition of corrective material, should be favored over quiet retroactive alteration. And publishers should establish clear and principled policies against undertaking such changes under public pressure that falls short of a legal finding of infringement. (And, in plenty of cases, publishers should stand up against legal pressure, too.) The benefit of retroactive correction in some instances—imagine fixing a typographical error in the proportions of a recipe, or blocking out someone’s phone number shared for the purposes of harassment—should be contextualized against the prospect of systemic, chronic demands for revisions by aggrieved people or companies single-mindedly demanding changes that serve to eat away at the public record. The public’s interest in seeing what’s changed—or at least being aware that a change has been made and why—is as legitimate as it is diffuse. And because it’s diffuse, few people are naturally in a position to speak on its behalf.

For those times when censorship is deemed the right course, meticulous records should be kept of what has been changed. Those records should be available to the public, the way that Lumen’s records of copyright takedowns in Google search are, unless that very availability defeats the purpose of the elision. For example, to date, Google does not report to Lumen when it removes a negative entry in a web search about someone who has invoked Europe’s “right to be forgotten,” lest the public merely consult Lumen to see the very material that has been found under European law to be an undue drag on someone’s reputation (balanced against the public’s right to know).

In those cases, there should be a means of record-keeping that, while unavailable to the public in just a few clicks, should be available to researchers wanting to understand the dynamics of online censorship. John Bowers, Elaine Sedenberg, and I have described how that might work, suggesting that libraries can again serve as semi-closed archives of both public and private censorial actions online. We can build what the Germans used to call a *giftschrank*, a “poison cabinet” containing dangerous works that nonetheless should be preserved and accessible in certain circumstances. (Art imitates life: There is a “restricted section” in Harry Potter’s universe, and an aptly named “poison room” in the television adaptation of *The Magicians*.)

It is really tempting to cover for mistakes by pretending they never happened. Our technology now makes that alarmingly simple, and we should build in a little less efficiency, a little more inertia that previously provided for itself in ample qualities because of the nature of printed texts. Even the Supreme Court hasn't been above a few retroactive tweaks to inaccuracies in its edicts. As the law professor Jeffrey Fisher said after our colleague Richard Lazarus discovered changes, "In Supreme Court opinions, every word matters ... When they're changing the wording of opinions, they're basically rewriting the law."

On an immeasurably more modest scale, if this article has a mistake in it, we should all want an author's or editor's note at the bottom indicating where a correction has been applied and why, rather than that kind of quiet revision. (At least, I want that before I know just how embarrassing an error it might be, which is why we devise systems based on principle, rather than trying to navigate in the moment.) Society can't understand itself if it can't be honest with itself, and it can't be honest with itself if it can only live in the present moment. It's long overdue to affirm and enact the policies and technologies that will let us see where we've been, including and especially where we've erred, so we might have a coherent sense of where we are and where we want to go.