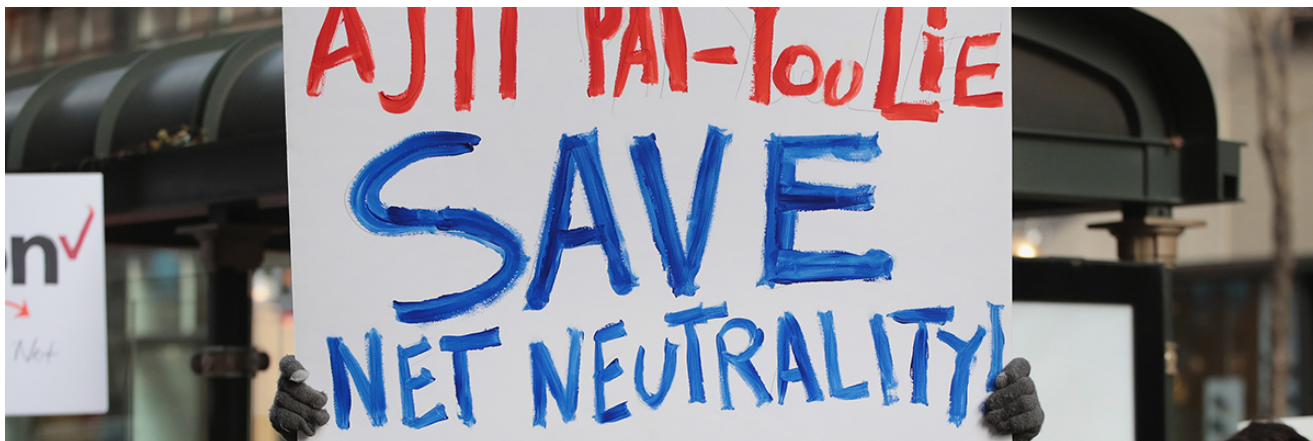


Amenazas contra la libertad de expresión en Internet

[Lino González Veiguela](#)



¿Cómo limitan la libertad de expresión en la Red los gobiernos y Estados del mundo, no sólo los autocráticos también las democracias? He aquí las claves para entenderlo y cuáles son los retos.

En su último [informe](#), la ONG Comité para la protección de los periodistas ofrecía una lista de los diez países con menos libertad informativa. Se trata de países en los que se lleva a cabo una combinación de tácticas de represión tradicionales –como el encarcelamiento de periodistas, la censura sistemática o el cierre de medios- con un uso progresivo de herramientas electrónicas que les permiten ejercer un control cada vez más efectivo de los contenidos publicados en Internet.

Las amenazas a día de hoy para la libertad de expresión en Internet, sin embargo, no proceden únicamente de la represión estatal en países autocráticos como Eritrea, Corea del Norte, Turkmenistán, Arabia Saudí o China, con su modélico *tecnoautoritarismo*. En las democracias, se ha ido conformando un ecosistema mediático en el que interactúan empresas tecnológicas que obtienen beneficios por el volumen de información compartida y el tiempo de atención que le dedican los usuarios; gobiernos y partidos políticos cada vez más activos a la hora de contribuir a campañas de desinformación y medios de comunicación que tratan de viralizar contenidos para detener la sangría de lectores y publicidad. Ese ecosistema enrarecido está afectando a los flujos informativos cotidianos y también, con más intensidad, a los que se producen durante las campañas electorales. Todo ello en un contexto de desigualdad creciente, precarización laboral y desconfianza generalizada entre los ciudadanos hacia sus

representantes públicos. La crisis provocada por la pandemia del coronavirus no ha hecho más que intensificar, en casi todos los países, dinámicas que ya estaban en curso desde mucho antes de que China emitiera una alerta sanitaria.

¿El fin de la neutralidad de la Red?

La libertad de expresión e información en Internet depende, en buena medida, de que podamos acceder a todos los contenidos a la misma velocidad, sea el vídeo de un gran medio de comunicación o uno incrustado en un blog con unos pocos cientos de lectores. La neutralidad en Internet implica que los proveedores de red –las grandes empresas de telecomunicaciones– no pueden acelerar el acceso de unos sitios web ni ralentizar la entrada a otros. En 2017, sin embargo, la Comisión Federal de Comunicaciones de Estados Unidos (FCC, por sus siglas en inglés) [trató de cambiar](#) la regulación para permitir que los grandes proveedores de conexión –como Verizon, Comcast o AT&T– pudieran negociar con los proveedores de servicios –como Netflix o con medios de comunicación con más recursos– un precio de conexión si querían que sus contenidos fueran accesibles preferentemente. Varios de estos proveedores de servicios, aliados con gobiernos estatales, presentaron un recurso contra esa decisión.

El pasado octubre un tribunal de apelación emitió una sentencia ambivalente: reconoció la potestad de la FCC para poner fin a la neutralidad de Internet, pero no aceptó su competencia para imponer esa regulación federal a los estados. Además de los recursos judiciales planteados contra esa sentencia, serán claves las próximas elecciones para el futuro de esta regulación: los demócratas [pidieron](#) en el Congreso la revocación de la nueva regulación de la FCC, pero no lograron que el Senado, controlado por los [republicanos](#), secundase la iniciativa.

En paralelo a esta lucha judicial y política, cada vez son más numerosas las voces que piden otro debate sobre las plataformas de contenidos y buscadores. Basándose en unos [algoritmos opacos](#), éstas obtienen una parte considerable de sus ingresos discriminando contenidos sin rendir cuentas. Una falta de neutralidad de la Red relacionada con la preeminencia de unos contenidos frente a otros que, de momento, los reguladores no se han atrevido a abordar. Las multas contra Google [impuestas por Bruselas](#), por ejemplo, han tenido que ver hasta ahora con una regulación antimonopolio basada en la concepción tradicional de la competencia, necesaria pero insuficiente para afrontar las nuevas fuentes de poder de los gigantes digitales.



Desconexión total de Internet

Durante 2019, las autoridades indias suspendieron Internet en casi 100 ocasiones. La mayoría de [desconexiones](#) duraron unas cuantas horas y se limitaron a estados concretos durante episodios de tensiones sociales (aunque en Jammu y [Cachemira](#), por ejemplo, no se restauró el tráfico de baja velocidad hasta el [pasado marzo](#), manteniéndose las restricciones sobre el 4G). También se produjeron suspensiones generalizadas en todo el país, como en diciembre, tras las protestas que siguieron al endurecimiento de los requisitos de ciudadanía para las minorías. India encabeza la lista de Estados que más cortes de la Red han llevado a cabo durante 2019 (también [en 2018 estuvo la primera](#): Pakistán fue el segundo país con más suspensiones, a mucha distancia de su vecino). En lo que va de 2020, el Gobierno indio ha llevado a cabo ya [6 desconexiones masivas](#) en diversos estados.

La [mayor democracia del mundo](#) –como suele denominarse– no es el único país que recurre a estos [cortes](#) con una finalidad de control social. Otros como Yemen, Pakistán, [Bangladesh](#), [Etiopía](#), [Birmania](#) o Irak han recurrido en los últimos meses a esta acción de censura, al menos, en algunas zonas de sus Estados. Los motivos alegados [varían](#), aunque casi siempre se justifican para evitar un mal mayor: impedir la violencia, reducir la inestabilidad política, combatir las noticias falsas o facilitar el desarrollo de un proceso electoral. La pandemia del coronavirus

ha sido la última excusa perfecta para muchos de ellos, según ha [denunciado](#) Human Rights Watch. Sin embargo, las organizaciones de derechos humanos afirman que el motivo principal de estas desconexiones es evitar que la información sobre la represión trascienda o lo haga con mayor lentitud y menor impacto. Los objetivos prioritarios de los cortes masivos de Internet son, sobre todo, las redes sociales. Otros países, [como Venezuela](#), se han centrado en cortes más selectivos, como declaraciones de opositores *por streaming*.

Una práctica que [comenzó](#) en China en 2009, como acción represiva en la provincia de Xinjiang, se ha ido extendiendo en estos últimos 10 años llevándose ya a cabo en más de treinta países. En 2016, se produjeron 75 cortes en todo el mundo, 106 en 2017 y 196 en 2018, [según](#) la ONG Access Now. Sólo en la primera mitad de 2019 se confirmaron 128 cortes. De estas cifras, se excluyen países como China y Corea del Norte, con una censura permanente y sistemática. Las autoridades de Pekín ven con buenos ojos esta tendencia: un artículo publicado en el *China Daily* –[órgano de propaganda](#) del régimen- [afirmaba](#) el pasado diciembre que todo Estado soberano debería tener derecho a cortar el flujo de Internet si lo consideraba oportuno para defender sus intereses. Soberanía digital al servicio de la represión. A finales de marzo, en plena crisis de salud pública mundial, las autoridades chinas anunciaron que presentarían [una propuesta](#) ante Naciones Unidas para reconfigurar la Red que permitiría una desconexión selectiva de contenidos y un control más eficaz del flujo informativo por parte de operadores y gobiernos.

El último gran episodio de corte masivo de Internet tuvo lugar en noviembre, cuando las autoridades [iraníes](#) fueron capaces de dejar en una casi total oscuridad digital a sus 81 millones de habitantes durante varios días. Ha sido, según los expertos, uno de los cortes más precisos registrados hasta la fecha. Todo un ejemplo para autoridades de otros Estados que pretendan seguir el ejemplo de las autoridades de Teherán para silenciar a sus poblaciones.

Además de suponer una vulneración de uno de los derechos humanos más básicos –la libertad de expresión y de información-, los cortes de Internet también tienen un importante coste económico. Algunas estimaciones [cifran](#) en 8.000 millones de dólares anuales las pérdidas que generan.



Feudalización de Internet

Según fuentes oficiales, las autoridades iraníes habrían aprovechado el gran corte de la Red de 2019 para comprobar el funcionamiento de su proyecto de Internet nacional, un “Internet halal” que impediría, de lograrse su implantación, cualquier acceso a contenidos no aprobados por las autoridades del régimen teocrático. El proyecto, según las [estimaciones](#) de algunos expertos, podría estar operativo en tres años, logrando un Internet totalitario y autosuficiente. El “intranet” iraní no es únicamente el proyecto del ala más dura del régimen: el primer ministro Rohaní [ha declarado](#) que esta red soberana permitirá que los iraníes puedan satisfacer sus “necesidades” de conexión sin recurrir a redes extranjeras.

Tan sólo unos días después del gran corte masivo en Irán, Rusia [anunció](#) que había completado satisfactoriamente una prueba de su futuro Internet soberano, conocido como “Runet”. Según las autoridades rusas, la prueba se llevó a cabo sin que los internautas rusos percibieran ninguna disrupción. La Red, según expertos consultados por medios rusos, podría estar [lista](#) para 2021. Aunque otros investigadores [ponen en duda](#) que el proyecto esté disponible en un plazo tan breve. También se podría optar por un modelo de censura similar al chino –con un gran cortafuegos controlado por el Gobierno-, en lugar de por una desconexión absoluta de la Red mundial.

Más allá de proyectos con una clara dimensión represiva, la “soberanía en Internet” – en especial, la que se da sobre los datos- preocupa a todos los países. También a la Unión Europea. Europa se ha dado cuenta, por fin, de que en el precio que se paga por tecnologías

de la información extranjeras se ha de incluir también el coste de la pérdida de la soberanía. Y ese coste no resulta fácil de cuantificar, considerando que la información es desde hace tiempo una de las diversas dimensiones esenciales de la soberanía de cualquier comunidad política.

Quién controla los datos

El gran poder de las grandes empresas tecnológicas, en términos de acumulación de datos cuyo uso no está sometido a escrutinio, fue objeto de debate en la [reciente reunión](#) entre el fundador de Facebook y la comisaria de Justicia de la Unión Europea. La consigna de Bruselas es que “es Facebook quien tiene que adaptarse a la regulación de la UE, no al revés”. Una consigna que, al menos en teoría, la Unión quiere mantener frente al resto de empresas tecnológicas. El último reglamento europeo sobre protección de datos puede considerarse, en este sentido, un primer aviso. La cuestión es si la UE podrá hacer respetar su criterio careciendo de infraestructuras tecnológicas propias comparables al músculo tecnológico de Estados Unidos y China. Además de la lógica preocupación por el gran retraso tecnológico europeo –con las implicaciones geopolíticas que ello comporta-, en Bruselas y en varias capitales europeas se quiere actuar en el ámbito de la gestión y el almacenamiento de datos por parte de las grandes empresas del sector, estadounidenses, en su mayoría. Esa es la motivación de Francia y Alemania para haber anunciado recientemente un [proyecto de almacenamiento de datos en la nube](#) que reduzca la dependencia de los grandes servidores norteamericanos. El proyecto tiene ya el [respaldo de la Comisión](#). Las consecuencias económicas y geopolíticas dependerán del grado de independencia respecto de los operadores estadounidenses.

El comisario europeo para el mercado interior, [Thierry Breton](#) –director hasta hace unos meses de la compañía tecnológica Atos- se muestra [optimista](#) sobre las posibilidades de futuro de la UE, afirmando que el 60% de las patentes relacionadas con el 5G son europeas. A largo plazo, esas patentes pueden ser la base para una tecnología europea. Por el momento, sin embargo, varios países –Alemania y Francia, entre ellos, así como España- ya han [anunciado](#) su intención de desplegar el 5G sirviéndose de la infraestructura china de Huawei. Al menos en las primeras fases del despliegue, a la espera de poder utilizar la tecnología de empresas europeas como Nokia o Ericsson. Estados Unidos se ha [manifestado](#) ya en contra de esta alianza con la empresa china. Europa toma nota de la queja, pero el precio ofrecido por Huawei es más barato.

Algoritmos opacos

A la espera de conocer las repercusiones de esta asociación con Huawei, en los países democráticos las mayores amenazas contra la libertad de expresión en Internet provienen, sobre todo, de las propias compañías tecnológicas occidentales. Además de [colaborar](#) activamente y de modo opaco con varios servicios secretos occidentales, grandes compañías como Alphabet están favoreciendo, según diversos estudios, que las tendencias informativas de los internautas con sesgos extremistas se vean alimentadas con noticias cada vez más extremistas. Muchos usuarios estarían recibiendo contenidos que profundizan su radicalización conforme van accediendo a los vídeos propuestos por la plataforma en función de sus primeras búsquedas. En otras palabras, compañías como Youtube ofrecen al consumidor un menú de contenidos que [maximiza la economía de la atención](#) y logra, por tanto, un aumento de ingresos publicitarios -13.000 millones de euros [en 2019 declarados por la plataforma de vídeos](#)— obtenidos gracias a la satisfacción compulsiva de las demandas de los usuarios que viven en *cámaras de eco* informativo.

Un [estudio](#) sobre Youtube presentado en enero 2020, detectó que entre los usuarios se podía identificar una tendencia del algoritmo a ofrecer contenidos cada vez más extremistas, como los suscriptores que buscaban vídeos de extrema derecha. La compañía ha cuestionado las conclusiones de esa investigación académica. Los [autores de la investigación](#) critican, a su vez, la negativa de la empresa a ofrecer acceso al funcionamiento de sus algoritmos. Las conclusiones del estudio son [consistentes](#) con otras investigaciones realizadas en los últimos años sobre diversas redes sociales. En Twitter la amenaza [proviene](#) de la querencia del algoritmo por favorecer la viralidad de las noticias falsas.

Facebook, también señalado en este sentido, ha [comenzado](#) ya a aplicar una política de detección de noticias falsas, [subcontratando](#) este servicio a diversas compañías y medios en varios países. Incluida [España](#), donde se ha aliado con la agencia de noticias AFP, con Newtral y con Maldita, lo que motivó un [debate público](#) durante las primeras semanas de la crisis sobre el alcance de esta colaboración. Pero esta no parece que sea la solución definitiva, ni la más adecuada para garantizar la libertad de expresión.

De las *fake news* al *deep fake*

Las campañas electorales del *Brexit* y de las últimas elecciones presidenciales de Estados Unidos y Brasil (con Trump y un hijo de Bolsonaro [investigados](#)) alertaron sobre la expansión de las noticias políticas falsas. Desde entonces se está debatiendo sobre qué [medidas](#) pueden

tomarse al respecto.

Numerosos países han aprobado leyes –o las están tramitando- para afrontar el problema. En [Singapur](#) ya se puede condenar a una persona a más de 10 años de cárcel por difundir noticias falsas o información que el Gobierno considere propaganda antigubernamental, y las empresas que las difundan también podrán ser multadas. Francia [ha aprobado una ley](#) que permite a los jueces eliminar noticias falsas durante las campañas electorales. Además, varios países cuentan ya con unidades operativas para “combatir las *fake news*”, como las [creadas](#) en España –en línea con las de la Unión Europea- durante las pasadas campañas electorales.

El Gobierno australiano quiere ir más allá y pretende [aprobar](#) una ley que permita a una agencia administrativa filtrar los contenidos que considere inapropiados sin intervención judicial. El Ejecutivo británico también ha presentado un [plan](#) de control de contenidos en Internet que ha generado polémica, aunque aún no ha entrado en los cauces parlamentarios. En India, las autoridades han presentado un [borrador de ley](#) que les facultaría para pedir a las compañías que suspendan Internet durante 24 horas. La ley [aprobada a finales de marzo](#) en Hungría a propuesta de Viktor Orban para afrontar la pandemia, ya habría provocado las primeras detenciones de personas por publicar en redes sociales informaciones que se han demostrado ciertas sobre capacidades hospitalarias o [comentarios](#) en los que se llamaba dictador al presidente (los cargos fueron finalmente retirados).

Al mismo tiempo, Facebook –propietaria de Whatsapp-, que ha creado la ya mencionada unidad centrada en “combatir las *fake news*”, también ha impuesto un [límite](#) al número de veces que se puede reenviar un mensaje a grupos de contactos, para limitar la potencial expansión de noticias falsas. En las últimas semanas, [ha borrado](#) contenidos que contenían –según la compañía- [falsedades](#) sobre la pandemia (aunque algunos medios de comunicación denunciaron que les habían eliminado contenidos veraces, incluidos [medios españoles](#)). A pesar de todo ello, un [estudio de la fundación Avaaz](#) ha señalado que la compañía está siendo incapaz de detectar el 70% de los bulos en español que se publican en su red.



Por su parte, Twitter acaba de anunciar que [señalará](#) la información que considere dudosa o falsa sobre el coronavirus, conminando a buscar “los hechos reales” sobre la enfermedad. En abril, ya había [anunciado](#) que suprimiría contenidos relacionados con la pandemia que supusiesen “un riesgo para la salud”.

Todas estas medidas pueden contribuir a limitar el impacto de la desinformación, pero plantean [un debate](#) sobre si han de ser las compañías tecnológicas las que se encarguen de una labor que, en muchos casos, comporta aplicar serias restricciones a la libertad de expresión. Les convierte en los editores –que dicen no ser, para evitar someterse a las reglas de los medios de comunicación tradicionales- y en jueces –que, desde luego, ni son ni deberían ser-.

La cuestión a nivel regulatorio es compleja. Uno de los grandes desafíos, si queremos evitar que la libertad de expresión quede en manos de compañías privadas o en decisiones administrativas, será actualizar -y acelerar- el modo de intervención de los jueces en los procesos de “censura preventiva”. El problema es que la velocidad de los procesos y el nivel de garantías de los procesos son, en muchos sentidos, variables inversamente proporcionales.

Otro problema añadido a la dificultad regulatoria, tiene que ver con el uso opaco de las redes que están haciendo casi todos los partidos políticos –con uso de *bots*, por ejemplo- para ejercer

la oposición o justificar sus medidas del Gobierno. Mientras obtengan beneficios demoscópicos y electorales de ese comportamiento, los alicientes para limitar con una ley adecuada su propio margen de actuación son escasos. En España, por ejemplo, se supo en septiembre que Facebook y Twitter habían cerrado unas 350 cuentas creadas con identidades falsas vinculadas –[según las compañías](#)– al Partido Popular. Fueron usadas para publicar *spam* durante la campaña electoral de abril de 2019, sin fiscalización de las Juntas electorales –otros organismos claramente desfasados respecto al ámbito digital-. Un estudio de la Universidad de Murcia [ha vinculado](#) unos 41.000 bots activos durante la última campaña electoral de noviembre con una intensa actividad amplificando mensajes en apoyo de los principales partidos políticos españoles. El estudio no tenía entre sus objetivos probar la vinculación de esos *bots* con las formaciones políticas. En abril, Facebook [borró cuentas falsas](#) que habían saturado el perfil institucional del Ministerio de Sanidad español, relacionadas –según la compañía- con una red de *spam* global.

Esas dificultades regulatorias aumentarán cuando se perfeccionen y extiendan las imágenes denominadas “*deep fake*”. Las nuevas herramientas de inteligencia artificial están permitiendo crear vídeos en los que personas reales hacen o dicen cosas que nunca dijeron o hicieron. La calidad de estos vídeos es cada vez mayor, y terminará resultando imposible distinguir entre uno real o falso.

Hace unos meses, tras la negativa de Facebook a retirar un vídeo *deep fake* de Nancy Pelosi en la que ‘aparecía’ borracha, unos artistas subieron otro vídeo de Mark Zuckerberg en el que le [hacían presumir](#), entre otras cosas, de que robaba datos de millones de personas. La tecnología de los *deep fake* [podría llevar](#) la desinformación a otro nivel, especialmente durante los breves períodos de campañas electorales o de crisis como la pandemia actual. Aunque también hay expertos que aseguran que su impacto político será mucho [menor](#) del anunciado. Entre otras cosas, porque herramientas que usen también la inteligencia artificial podrían [detectar fácilmente](#) los *deep fake*, utilizados hasta la fecha sobre todo en la industria del porno.

El problema con los vídeos *fakes* políticos – que comparten con las noticias falsas y los bulos- es que no tienen como objetivo prioritario imponerse a la verdad para sustituirla permanentemente: les basta con distorsionarla de forma temporal. Si logran que una parte considerable del público crea en su contenido –durante unas horas o unos días como máximo-, habrán tenido éxito, sin importar demasiado que se practique una verificación posterior e incluso que se proceda a su retirada de la Red.

En las pasadas elecciones estatales de Nueva Delhi, el jefe del partido de Modi en el estado –un exactor de Bollywood-, subió a las redes un vídeo en el que aparecía hablando con soltura

en un dialecto que apenas conoce para acusar al Gobierno en funciones de mentir. Su intención era hacer campaña entre el electorado que sí habla ese dialecto. El vídeo se subió a redes justo un día antes de las elecciones. [Según los responsables](#) de campaña del candidato, este contenido falso circuló por unos 5.800 grupos de WhatsApp y podría haber llegado a unos 15 millones de personas –más que el electorado potencial-. Varios medios indios y occidentales –incluidos algunos medios españoles- acusaron al político de haber usado tecnología de *deep fake* para lograr votos.

Al final, [se demostró](#) que no era un sofisticado vídeo *deep fake*, sino uno simple manipulado: lo único falso era la voz del político, quien, sin embargo, demostró un gran talento para sincronizar sus labios con el mensaje. La buena noticia es que su partido, a pesar de utilizar una estrategia engañosa, no logró la victoria. La mala: que consiguieron aumentar sus votos y –casi- triplicar sus escaños, pasando de 3 a 8 ([muy lejos](#), en todo caso, de la victoria). Resulta imposible saber cuántos votos ganó con ese vídeo, pero a su partido no le han retirado, hasta la fecha, ninguno de los escaños obtenidos.

Fecha de creación

14 mayo, 2020