



La Trama de la Comunicación

ISSN: 1668-5628

ISSN: 2314-2634

latramaunr@gmail.com

Universidad Nacional de Rosario

Argentina

Zelcer, Mariano

Machine learning y lógicas semióticas: el caso de la publicidad digital
La Trama de la Comunicación, vol. 26, núm. 2, 2022, Julio-Diciembre, pp. 15-31
Universidad Nacional de Rosario
Rosario, Argentina

Disponible en: <https://www.redalyc.org/articulo.oa?id=323973878001>

- Cómo citar el artículo
- Número completo
- Más información del artículo
- Página de la revista en redalyc.org

redalyc.org

Sistema de Información Científica Redalyc

Red de Revistas Científicas de América Latina y el Caribe, España y Portugal
Proyecto académico sin fines de lucro, desarrollado bajo la iniciativa de acceso
abierto

Machine learning y lógicas semióticas: el caso de la publicidad digital

Por Mariano Zelcer

marianozelcer@yahoo.com.ar - IIEAC (Instituto de Investigación y Experimentación en Arte y Crítica),
Área Transdepartamental de Crítica de Artes Oscar Traversa, Universidad Nacional de las Artes (UNA), Argentina

RESUMEN:

Este artículo propone una aproximación a los procesos de aprendizaje automático por computadora desde una perspectiva semiótica peirciana. Para ello, trabaja en un territorio privilegiado para la observación de la articulación de la dataficación de los usuarios con su posterior gestión mediante sistemas informáticos de aprendizaje por computadora: la publicidad digital. A partir de un caso real, se da cuenta de los modos en los que el machine learning articula lógicas abductivas e inductivas, poniendo foco en los modos en que los sistemas informáticos generan hipótesis a partir de la identificación de semejanzas y las ponen a prueba en investigaciones experimentales, cuyos resultados funcionan como input que realimenta el aprendizaje.

DESCRIPTORES:

publicidad digital, *machine learning*, dataficación, semiótica, abducción

ABSTRACT:

This article introduces an approach to the computer machine learning processes from Peirce's semiotics perspective. To reach this aim, it works in a territory where the machine learning management of data obtained through users' datafication becomes clearly visible: digital advertising. Based on a real case, it describes the ways in which machine learning articulates abductive and inductive logics, focusing on how computer systems generate hypotheses based on the identification of similarities and test them in experimental research, whose results work as a new input that feeds back into learning.

DESCRIPTORS:

digital advertising, machine learning, datafication, semiotics, abduction

15

Machine learning y lógicas semióticas: el caso de la publicidad digital

Machine learning and semiotic logics: the case of digital advertising

Páginas 015 a 031 en La Trama de la Comunicación, Volumen 26 Número 2, julio a diciembre de 2022

ISSN 1668-5628 - ISSN 2314-2634 (en línea)

1. POR QUÉ ESTUDIAR LA PUBLICIDAD EN INTERNET¹

La progresiva digitalización de las distintas actividades cotidianas, en múltiples esferas de la acción social, ha implicado —entre otras cuestiones— que se generen, en forma permanente, datos acerca de entidades de muy distinta naturaleza en su existencia física, en un fenómeno que se conoce como *datificación*. Martín Gendler la explica que con la datificación se produce una “cuantificación y trazabilidad de cada instancia de la vida social a través de la sistematización algorítmica de los mecanismos de obtención, procesamiento, aplicación y puesta en juego de datos de los sujetos, objetos y prácticas de los usuarios” (Gendler, 2021:21).

La enumeración que propone Gendler (sujetos, objetos y prácticas) da cuenta de tres grandes universos de datos que, en verdad, se encuentran relacionados: los estudios que se han desplegado en los últimos años acerca de los sistemas de recomendación de las plataformas de contenidos digitales² han dado un ejemplo de ello. En muchos casos, entre los que se encuentran esos entornos en los que se producen consumos de productos culturales, las prácticas implican la vinculación de sujetos con objetos: cada escucha de un tema musical en Spotify o visualización de una película en Netflix genera datos referentes a una práctica, que luego funcionan como *input*³ en el

mundo de los algoritmos para la generación de una nueva recomendación.

Es sabido que los datos recolectados mediante estos sistemas no se emplean únicamente para brindar recomendaciones o experiencias personalizadas, sino también con múltiples otros fines. Entre ellos, uno de los centrales es la distribución de publicidad. La publicidad es, para muchas plataformas y conglomerados de plataformas, el corazón del negocio con el cual sostienen su funcionamiento y generan ganancias, muchas veces extraordinarias.⁴ Se puede pensar que la datificación está, en esos entornos, programada para funcionar de modo eficiente en dos frentes. Por un lado, el propio de la plataforma: así, YouTube buscará —por ejemplo— ser lo más relevante posible en los videos que sugiere al usuario o en las respuestas que da cuando se realizan búsquedas; Facebook buscará ordenar y mostrar, de entre las publicaciones de los contactos y marcas que sigue cada usuario, aquellas que le resulten más relevantes en primera posición, y algo similar realizará Instagram, ordenando y proponiendo a cada usuario contenidos personalizados. Por el otro, la datificación está, en buena medida, puesta

de aquello que constituye una entrada para un algoritmo informático. Como recuerda Raquel Maluenda de Vega, todos los algoritmos informáticos poseen tres partes: 1) la entrada o *input*, con la que el algoritmo va a trabajar para alcanzar la solución que se espera; 2) el proceso, que es específicamente el conjunto de pasos o instrucciones para llegar, a partir de la entrada, a la salida, también conocida como *output*, y 3) un *output* o salida, que presenta los resultados (Maluenda de Vega, 2021).

4 De hecho, varias de las empresas que manejan las principales plataformas se encuentran entre las más valiosas del mundo. Atendiendo a la capitalización del mercado, Alphabet Inc. (empresa propietaria de Google) y Meta Platforms Inc. (propietaria de Facebook y de Instagram) se encontraban, a comienzos de 2022, entre las 10 de mayor valor (ver <https://es.fxssi.com/las-empresas-mas-valiosas-del-mundo>, consulta 19/2/2022)

1 Este artículo ha sido elaborado en el marco del proyecto PIACyT 34/0585 “Propuestas curatoriales y lógicas automatizadas en plataformas de contenidos: una caracterización semiótica de los sistemas de recomendación que operan por *machine learning*”, radicado en el Instituto de Investigación y Experimentación de Arte y Crítica (IIEAC) del Área Transdepartamental de Crítica de Artes Oscar Traversa de la Universidad Nacional de las Artes, que se lleva a cabo en el marco de la programación científica 2020-2021 (extendida hasta 2022).

2 Ver, por ejemplo, Ricci, Rokach y Bracha, 2015; Cingolani, 2017a y 2017b o Koldobsky, 2019.

3 Empleamos el término “input” en el sentido específico

en función de su posterior empleo con fines publicitarios: cuanto más conoce una plataforma acerca de cada usuario, sus comportamientos y sus preferencias, podrá mostrarle publicidad que le resulte relevante, y así, sostener y desarrollar su negocio.

Al analizar el funcionamiento de las soluciones publicitarias asociadas a las plataformas, accedemos entonces a un punto de observación privilegiado en el que se hacen visibles muchos aspectos referentes a la datificación, que no se pueden observar, o son menos evidentes, cuando se analizan las interfaces (Scolari, 2018) con las que interactúan los usuarios. El mundo de la publicidad digital, además, permite poner el foco con mucha mayor atención en los modos en que la datificación opera con los sujetos, suspendiendo en buena medida la pregunta por su interacción con los objetos, que se produce en el otro frente.

2. GESTIÓN VERSUS RECOLECCIÓN

Históricamente, las soluciones de publicidad en las grandes plataformas digitales como Facebook (hoy Meta) y Google requerían la definición, por parte de quien implementaba las campañas, de múltiples variables. Entre ellas, muchas se referían a las características de los usuarios que se deseaba alcanzar. En el mundo de la publicidad, esto se conoce como "segmentación". Una observación de las interfaces de las plataformas publicitarias permitía identificar las distintas variables con las cuales cada plataforma permitía seleccionar los usuarios a los cuales se les orientarían las piezas publicitarias. Esta manera de hacer publicidad sigue estando vigente y, de hecho, las herramientas de segmentación son un espacio privilegiado para observar la datificación, en la medida en que hacen visible buena parte de los datos que las plataformas tienen asociados a los usuarios, que van desde algunos de los llamados "duros", más ligados a características sociodemográficas (sexo, edad, lugar de residencia, etc.), hasta otros que dan cuenta de sus

intereses o comportamientos.⁵ En un funcionamiento de este tipo, las plataformas *recolectan* los datos de los usuarios, luego el operador define los parámetros o la segmentación que empleará en sus campañas, y finalmente la plataforma orienta la publicidad al público o "segmento" programado.⁶

En los últimos años, sin embargo, las soluciones publicitarias de estas plataformas fueron desarrollándose y orientándose, cada vez más, a mecanismos "inteligentes" de selección de audiencias. En estos casos, el operador sólo indica unos pocos datos, y la campaña se ejecuta con métodos de aprendizaje automatizado; es decir, la plataforma aprende a través de la experiencia, y entiende, sin que un operador se lo indique, a quiénes debe mostrarles los anuncios para lograr los objetivos deseados en una campaña. En estos casos, en los que intervienen algoritmos de *machine learning*, las plataformas no sólo recolectan datos, sino que también *gestionan* en buena medida la campaña, con menor intervención de los operadores.

3. EL MACHINE LEARNING

El *machine learning* (que se ha traducido alternativamente como "aprendizaje automático", "aprendizaje automatizado" o "aprendizaje de máquinas") es una rama de la inteligencia artificial cuyo objetivo es desarrollar técnicas que permitan que las computa-

5 Por las políticas de resguardo de la privacidad, las plataformas por lo general no permiten saber quiénes son los usuarios que componen un grupo de una determinada segmentación, sino solamente cuántos son. Es decir, se pueden programar campañas orientadas a ellos, pero sin conocer sus identidades individuales. Estas identidades sólo se podrán conocer si cada usuario desea compartir la suya activamente (por ejemplo, si hace clic en un anuncio, llega a un formulario que le pide sus datos personales, los completa y los envía).

6 Nos hemos dedicado a estudiar este fenómeno en Zelcer, 2019.

doras “aprendan”. En este campo, se entiende por “aprendizaje” la capacidad de un agente informático (programa, aplicación) de mejorar con la experiencia; una mejora que puede ser evaluada tanto en relación con el cumplimiento de parámetros o reglas (por ejemplo, el empleo pertinente de las reglas de la lengua, en una plataforma de traducción automática) como por respuestas positivas de los usuarios de un sistema (por ejemplo, la visualización efectiva, por parte de los usuarios, de los productos audiovisuales seleccionados por un sistema de recomendación en un sistema de *streaming* de video). Por su naturaleza informática, el *machine learning* es considerado un subcampo de la inteligencia artificial, ubicado en el territorio más amplio de las ciencias de la computación. Junto al *machine learning*, en el campo de la inteligencia artificial se encuentran también los sistemas de razonamiento, los sistemas de planificación automática y los sistemas de procesamiento de lenguaje natural (ver Hurwitz y Kirsch, 2018:13); aunque estos últimos son incluidos también dentro de los sistemas de *machine learning* por varios autores.

El *machine learning* trabaja mediante algoritmos, que son secuencias de instrucciones o reglas definidas y ordenadas que permiten, mediante una serie de pasos lógicos, solucionar un problema, o más genéricamente, brindar una respuesta.

La expansión del *machine learning* a las plataformas y sistemas informáticos de empleo cotidiano ha sido notable en los últimos años. Son algoritmos de *machine learning* los que, por ejemplo, agrupan las fotografías en nuestros teléfonos celulares atendiendo a quiénes están fotografiados en ellas o traducen textos de un idioma a otro automáticamente en Google Translate; estos algoritmos deciden también qué publicaciones mostrarnos cuando ingresamos a plataformas como Facebook o Instagram, nos sugieren qué obra musical escuchar en Spotify, qué videos ver en YouTube o

qué serie podría interesarnos en Netflix.⁷ Como se ve en estos ejemplos, cuando los algoritmos de *machine learning* operan en el mundo de las interfaces de los usuarios, trabajan con los tres grandes universos señalados para la datificación (sujetos, objetos, prácticas). Sin embargo, cuando los sistemas de publicidad de las plataformas seleccionan automáticamente los usuarios a los que mostrarles los anuncios, estos algoritmos operan también articuladamente con el otro “frente”. Como se ve, a través de la datificación y el *machine learning*, las plataformas no sólo clasifican y organizan los productos culturales contenidos en ellas, sino que también lo hacen con los usuarios.

4. UN CASO: LOS NUEVOS CLIENTES DE UN BANCO

Para aproximarnos al estudio del *machine learning* en el caso de las plataformas publicitarias digitales, trabajaremos a partir de un caso real de publicidad en Internet, aunque cambiamos el nombre y el país del anunciante por una cuestión de confidencialidad. El Banco Solís, de Uruguay, nació hace unos años con un formato novedoso: sin sucursales, funciona únicamente en una app para celulares. Con el fin de ganar clientes, realiza campañas de adquisición de nuevos usuarios en las principales plataformas publicitarias de Internet (Google, Facebook/Meta). Ambas plataformas gestionan estas campañas con algoritmos de *machine learning*. Haremos foco ahora en Google, que es la plataforma que, al día en el que se escriben estas líneas, tiene más avanzado este tipo de funcionamiento, aunque, en términos generales, es similar en ambas.

7 Aunque no únicamente: en estas plataformas se dan procesos complejos para la selección y recomendación de contenidos, en los que se combinan procesos de selección editorial o manual con otros algorítmicos. Entre los algorítmicos, además, no todos son del tipo del *machine learning*, pero sí muchos de ellos. Un buen relevamiento para el caso de Spotify puede verse en Bonini y Gandini, 2019.

Una agencia de publicidad prepara diversas piezas publicitarias promocionando el banco e invitando a los usuarios a descargar la app en su celular: imágenes estáticas, otras con animaciones sencillas y algunos videos de corta duración en distintos formatos (con animaciones, banda musical y locución). A esto le suma algunos anuncios de solo texto.

Una vez listas las piezas, se programa la plataforma publicitaria, en un proceso que se llama "implementación". Como adelantábamos, la presencia del *machine learning* hace que el operador de la plataforma no deba indicar, en el momento de implementar la campaña, el segmento al que decide orientarla, sino que dejará esta tarea en manos de los algoritmos. Por ello, la programación de la plataforma publicitaria es relativamente sencilla.

Los pasos que sigue el operador pueden resumirse como sigue:

- Se cargan todas las piezas publicitarias
- Se indica el país en el que debe correr la campaña. Fuera de este dato, no se brinda ninguna otra indicación acerca de los usuarios que la campaña debe hallar (sexo, edad, intereses, nivel socioeconómico, profesión, etc.) ni acerca de dónde se desea que aparezcan los anuncios.
- Se determina el evento que se quiere que los usuarios realicen. En este caso, el evento consiste en la apertura de una cuenta en el Banco Solís, lo que implica que los usuarios descarguen la app, completen algunos datos personales, saquen una foto de su documento y graben un pequeño video de ellos mismos, todo exclusivamente con el celular. La app del banco está programada de modo tal que, cuando el usuario efectivamente termina exitosamente el registro (lo que equivale, en este caso, a abrir su cuenta bancaria), le envía una señal a la plataforma publicitaria que avisa que el objetivo se ha logrado. En el mundo de la publicidad digital, esto es lo que se

considera una "conversión".⁸

- Se define el monto máximo aproximado que se está dispuesto a pagar por cada nuevo cliente, es decir, por cada conversión.
- Y se determina un presupuesto máximo que se desea gastar por día.

La campaña comienza a correr, y les muestra los anuncios a usuarios cualesquiera en distintos espacios a los que la red de Google Ads tiene acceso (YouTube, Google Search, Google Play, múltiples sitios web y blogs, etc.). Algunos usuarios hacen clic, algunos de ellos descargan la app, unos pocos comienzan el registro y lo abandonan, y una parte muy pequeña completa todo el proceso y realiza una conversión. El algoritmo va mostrando los distintos anuncios, los somete a prueba, y, progresivamente, muestra preferiblemente aquellos que ayudan a "convertir" usuarios. A medida que se generan conversiones, el algoritmo "aprende": comienza a descubrir y a definir el perfil de los usuarios más propensos a descargar la app del banco y abrir una cuenta, y, en forma progresiva, exhibe los anuncios preferentemente a usuarios de ese perfil. Por un tiempo, el algoritmo permanece "en fase de aprendizaje", es decir, aún sigue aprendiendo acerca de los usuarios que pueden ser más propensos a descargar la app y convertirse en clientes. En forma estimada, al algoritmo le basta con conseguir 50 conversiones en una semana para finalizar este proceso. Algunos datos hacen visible este proceso: por ejemplo, la tasa de usuarios que abren una cuenta en el banco, comparada con los que sólo descargan la app sin finalizar esta operación, se va incrementando

⁸ En rigor, en el caso de Google Ads, la plataforma recibe también señales de otras acciones previas realizadas por el usuario; por ejemplo, la descarga de la app, y el llenado del primer formulario. Esto le permite "comprender" qué pasos va realizando cada usuario y, en el caso de aquellos que no finalizan el proceso, en qué punto lo abandonan.

con el tiempo, lo que indica que cada vez es mayor la proporción de usuarios “pertinentes” encontrados por el algoritmo. En muchos casos, además, el costo real de obtener cada uno de los nuevos clientes desciende, lo que indica que el funcionamiento del algoritmo se vuelve más eficiente.

Una vez finalizada la etapa de aprendizaje, la campaña sigue corriendo. Si el aprendizaje ha sido exitoso, cada día el banco conseguirá nuevos clientes por un costo igual o menor al que estaba dispuesto a pagar, y consumiendo como máximo el presupuesto diario que tenía asignado. Con el correr del tiempo, el algoritmo puede ir haciendo ajustes graduales, es decir, “continúa aprendiendo”, y múltiples circunstancias pueden hacer que la campaña tenga inconvenientes para correr, lo que probablemente lleve a un ajuste de alguna de las variables antes reseñadas.

¿Cómo funcionan aquí los algoritmos de *machine learning*? En principio, podemos decir que éste parece ser un caso de los llamados “aprendizaje no supervisado”:⁹ el algoritmo no es entrenado con

9 Los algoritmos de *aprendizaje supervisado* se caracterizan por establecer una correspondencia entre las entradas y las salidas deseadas del sistema; son típicos de este tipo de algoritmo los problemas de clasificación, en los cuales el sistema busca etiquetar (clasificar) un elemento escogiendo una entre varias categorías o clases. Por lo general, el aprendizaje supervisado comienza con una serie de datos en los cuales ya hay elementos clasificados, es decir que se le brindan al sistema, en un inicio, múltiples casos en los que hay tanto *inputs* como *outputs*. Hurwitz y Kirsch, explican que el aprendizaje, en este caso, “busca encontrar patrones entre esos datos que puedan ser empleados en el proceso analítico” (2018:15). En cambio, los algoritmos de *aprendizaje no supervisado* tienen la particularidad de llevar adelante sus procesos a partir de un conjunto de casos sólo por entradas al sistema –es decir, sin que se le proporcione al sistema información sobre las categorías en las cuales pueden clasificarse esos casos–; en estos tipos de algoritmo, el sistema debe tener la capacidad de reconocer patrones para poder etiquetar/clasificar las nuevas entradas, y eventualmente

ninguna selección muestral inicial previamente a su puesta en marcha, sino que aprende a partir de la experiencia. Cada nueva cuenta creada es un “evento” o señal que retroalimenta al algoritmo, es decir, vuelve a funcionar como *input*, y así, el proceso de aprendizaje avanza. Bazzara afirma que estos funcionamientos basados en la retroalimentación o el *feedback* fueron adelantados por la teoría cibernética, que brinda “la sustentación teórica sobre la cual (...) postular que es posible la comunicación entre humanos y máquinas, y entre máquinas y máquinas, y con ello (...) [sentó] las bases para el desarrollo, unos años después, de la inteligencia artificial (IA), de la cual se desprenderá –entre otras subdisciplinas– el Machine Learning” (Bazzara, 2021:42). Como señala este mismo autor, aquí hay un proceso de datificación a través del cual las actividades de los usuarios son “convertidas en datos que serán procesados por algoritmos” (Bazzara, 2021:42).

No podemos saber con exactitud qué es lo que ocurre dentro de los algoritmos, ni tampoco cuál o cuáles son los tipos específicos de algoritmos de *machine learning* empleados, pero, a partir de la observación de los resultados, sí podemos imaginar algunas cuestiones referentes al funcionamiento, que reseñamos a continuación.

La plataforma de publicidad (en este caso Google Ads, aunque Facebook/Meta funciona de modo muy semejante, y lo hacen de modo parecido muchas otras) dispone de copiosa información sobre los distintos usuarios: tanto datos “duros” sociodemográficos (como sexo, edad, nacionalidad, etc.), como otros

ajustar sobre la marcha el mismo sistema clasificatorio. Wilmott lo explica de este modo: “El aprendizaje no supervisado tiene lugar cuando los datos no están etiquetados. Esto significa que solo tenemos entradas [*inputs*], no salidas [*outputs*]; de algún modo, el algoritmo trabaja por sí solo” (Wilmott:2019:6). En todas las citas de Hurwitz y Kirsch y de Wilmott, la traducción es nuestra.

más “blandos”, referentes a intereses o intenciones: gustos, aficiones, comportamientos, etc. Las plataformas obtienen estos datos tanto a partir de información declarativa (los usuarios los han suministrado) como de otra que las plataformas generan a partir de sus comportamientos (qué sitios han visitado, a qué marcas o influenciadores siguen en redes sociales, a qué publicaciones les han dado “like”, qué han buscado en el buscador de Google, qué videos han visto en YouTube, a qué canales están suscriptos y muchísima información más).

Como dijimos, podemos imaginar que la campaña de publicidad del banco comienza mostrándoles anuncios a usuarios cualesquiera. La mayor parte de los usuarios los dejará pasar sin prestarles mayor atención; sin embargo, algunos los cliquearán; de entre los que cliqueen, algunos descargarán la app en su celular, y entre estos últimos, sólo algunos completarán todo el proceso de registro. En estos distintos pasos, se va produciendo una suerte de “embudo”, que ha sido largamente modelizado y estudiado en el ámbito del *marketing* digital. Cuando un primer usuario finaliza el registro y se convierte en un cliente, se produce una retroalimentación hacia el sistema.¹⁰ El funcionamiento se puede pensar, en términos de los tipos de algoritmo de *machine learning*, como de “aprendizaje por refuerzo”. En estos casos, el algoritmo aprende a través de una dinámica de ensayo y error: su información de entrada es el *feedback* o retroalimentación

que obtiene del mundo exterior como respuesta a sus acciones. En este tipo de aprendizaje, el agente informático debe aprender cómo se comporta el entorno mediante recompensas (llamadas “refuerzos”; de allí el nombre de este tipo de algoritmo) o castigos, derivados del éxito o del fracaso respectivamente. Hurwitz y Kirsch lo explican de este modo:

El aprendizaje por refuerzo es un modelo de aprendizaje conductual. El algoritmo recibe retroalimentación del análisis de los datos para guiar al usuario hacia el mejor resultado. (...) el sistema aprende mediante prueba y error. Por lo tanto, una secuencia de decisiones acertadas dará como resultado que el proceso se “refuerce” porque resuelve mejor el problema en cuestión (Hurwitz y Kirsch, 2018:16)

En estos casos, el *input* principal que el operador o instructor debe darle al sistema son los eventos que serán considerados éxitos o fracasos: “no hay un libro con reglas al cual referirse. Esto significa que el algoritmo tiene que aprender las consecuencias de tomar determinadas acciones del hecho mismo de tomarlas” (Wilmott, 2019:173).

Volvamos ahora a la campaña del Banco Solís. En este caso, la finalización del registro y la conversión de ese usuario en un cliente es el evento que se considerará un “éxito”, y que debería “reforzar” el aprendizaje del algoritmo. Imaginemos ahora que Google Ads dispone de abundante información acerca del primer usuario que ha abierto una cuenta en el banco: se trata de un varón, soltero, de 35 años, residente en la ciudad de Montevideo, hinchado del club Peñarol, amante de la tecnología, que usa Internet mayoritariamente en su teléfono celular, aficionado a ver *sitcoms* en las plataformas de *streaming*, a quien le gusta tanto la música clásica como el pop británico de la década de 1980. Evidentemente, cuando sólo hay un usuario que ha abierto la cuenta en el banco, el algoritmo no

¹⁰ Estamos explicando aquí el proceso en forma simplificada. Como adelantamos, cada una de las acciones que hace el usuario a lo largo del “embudo” puede funcionar como una señal, y entonces, el sistema de *machine learning* va aprendiendo no sólo cuando se genera el evento que indica que se ha conseguido un nuevo cliente, sino que lo hace a partir de todos los eventos que se van produciendo a lo largo del proceso. Sin embargo, empleando la terminología del *machine learning*, la “recompensa” de este algoritmo que aprende “por refuerzo” será la obtención de una nueva cuenta abierta en el banco. Volveremos sobre esto enseguida.

puede saber cuál o cuáles de estos datos son relevantes en términos de la construcción de un “perfil” de los usuarios que, con mayor probabilidad, descargarán la app del banco y abrirán su cuenta. La campaña continúa entonces mostrándoles los anuncios a muy diversos usuarios. Cuando varios usuarios ya hayan abierto su cuenta, el algoritmo podrá elaborar algunas hipótesis. Siguiendo con el ejemplo: si aparecen tanto varones como mujeres, el sexo no será una variable relevante. Si la edad de todos ellos, o de casi todos, es de 45 años o menos, entonces la campaña probablemente deba orientarse a un público que tenga, como máximo, esa edad. Si entre los usuarios hubiera varios aficionados al club Peñarol; esa podría ser una propiedad relevante; si los hubiera de otros clubes, entonces el gusto por el fútbol podría ser un interés relevante; si, en cambio, se mezclaran los hinchas de distintos clubes con otras personas que no tienen interés por ese deporte, entonces no sería una variable que el algoritmo deba considerar. Las principales plataformas publicitarias suelen señalar que requieren 50 conversiones para finalizar un proceso de aprendizaje.¹¹ Ese medio centenar de casos exitosos le permitirá al algoritmo entender cuál o cuáles son las propiedades relevantes de los usuarios que, con mayor probabilidad, realizarán una “conversión” (en el ejemplo del banco, la apertura de una cuenta), y cuáles no debe considerar.

En las campañas de publicidad digital que funcionan del modo tradicional, el operador que las implementa selecciona y define qué características, intereses o comportamientos deben tener las personas a las que se les muestran los anuncios. En ese caso, la lógica es de tipo deductivo: inicialmente se define la regla que determinará las audiencias (por ejemplo, “mujeres de 40 a 60 años residentes en Argentina con interés en

alimentación saludable”); luego, las plataformas decidirán si mostrar o no los anuncios a un usuario que se conecte, según sea un caso que sigue o no la regla. No hay aquí aprendizaje por parte del sistema; los criterios están definidos de antemano, y las plataformas simplemente “aplican una regla”.

En cambio, en las campañas que funcionan por *machine learning*, el criterio no está definido de antemano: es el mismo algoritmo el que irá determinando las características de los usuarios a los que mostrarles los anuncios a partir de los resultados que estos tengan en la interacción con sus anuncios. Si la lógica de las campañas tradicionales es deductiva, aquí entra en juego una *lógica abductiva* (Cheung, Steven, et. al., 2018).¹² La relación entre el funcionamiento de los algoritmos de *machine learning* y el tipo de razonamiento abductivo ya ha sido señalado por algunos estudios anteriores (ver, además de Cheung et al., Mehler et. Al, 2007, entre otros). Sin embargo, esos estudios –destinados a un público especializado en informática– realizan un señalamiento relativamente acotado de la presencia de esta lógica, y luego se dedican al desarrollo de fórmulas matemáticas que la vuelvan operativa en los funcionamientos informáticos, sin avanzar en una problematización que profundice la comprensión acerca de los modos en que estos razonamientos abductivos son realizados por los algoritmos desde la

23

12 Los autores hacen una aclaración muy relevante acerca del empleo del término “lógica”, que nosotros queremos hacer nuestra, y que es fundamental en un trabajo enmarcado en una investigación que se propone generar el diálogo entre distintas áreas del conocimiento: “el término ‘lógica’ no debe confundirse con ‘deducción lógica’. Lo estamos empleando en su sentido más amplio, el de cualquier sistema de reglas empleadas para llevar a cabo la investigación científica”. Lo mismo vale para cuando empleemos la expresión “tipos de razonamiento”: “razonar” tampoco será sólo realizar deducciones o inducciones, sino que incluirá también, como lo comprende Peirce (1965[1986]), los procesos abductivos. Volveremos sobre esto enseguida.

11 En algunos casos, incluyen también un límite temporal: por ejemplo, señalan que deben contabilizarse 50 casos en un lapso máximo de una semana.

lógica semiótica de Peirce. Nos proponemos realizar esa tarea en los apartados que siguen.

5. ABDUCCIÓN, INDICIALIDAD Y SEMEJANZA

La inducción, la deducción y la abducción son tres tipos de razonamiento que Peirce presenta en su obra en múltiples lugares. Cuando desarrolla sus tres tricotomías del signo, aparecen en la tercera de ellas, que es la que se ocupa de la relación del signo con su interpretante (rema, dicente y argumento) (ver Peirce, 1965[1986]:31). Estos tres tipos de razonamiento son presentados como tres tipos de argumento, que nosotros podemos pensar como tres tipos de lógica. Sin embargo, estos tres modos de razonar aparecen también en muchos otros pasajes de su obra, en los que los desarrolla y explica comparativamente. Debeamos reseñarlos, para luego reflexionar acerca del modo en el que pueden ayudarnos a comprender ciertos funcionamientos del *machine learning*, que se observan en el caso de la publicidad digital, pero que se extienden a otros espacios y plataformas.

La deducción y la inducción son relativamente conocidas; se estudian habitualmente en cursos de lógica y de pensamiento científico. La deducción trabaja con leyes o reglas que se consideran universales; de cualquier caso particular que corresponda a los propuestos por aquella ley se podrá llegar a una conclusión, que será una aplicación particular de esa regla general. Se trata del razonamiento que se corresponde con los silogismos lógicos, compuestos por una premisa mayor, una menor y una conclusión, como en el conocido ejemplo que dice "Todos los hombres son mortales" (premisa mayor) "Sócrates es hombre" (premisa menor) y "Sócrates es mortal" (conclusión). Peirce enfatiza que se trata de un tipo de razonamiento, es decir, de un modo de pensamiento, que no tiene que ver necesariamente con la realidad o el mundo exterior:

En la deducción, o razonamiento necesario, partimos

de un estado hipotético de cosas que definimos en ciertos aspectos abstractos. Entre los caracteres a los que no prestamos ninguna atención en este modo de argumento está el de si la hipótesis de nuestras premisas se conforma más o menos al estado de cosas en el mundo exterior o no. Consideramos este estado hipotético de cosas y somos llevados a concluir que sea como sea el universo en otros aspectos, sea donde sea y cuando sea que se realice la hipótesis, algo no explícitamente supuesto en esa hipótesis será invariablemente verdadero. (Peirce, 1998[2012b]: 278)

Por ende, y complementariamente, la conclusión será válida independientemente de ese mundo exterior:

Nuestra inferencia es válida si y sólo si hay realmente una relación tal entre el estado de cosas supuesto en las premisas y el estado de cosas afirmado en la conclusión. Si realmente es así o no, es una cuestión de la realidad y no tiene nada que ver en absoluto con cómo estemos inclinados a pensar. (Peirce, 1998[2012b]:278)

Suele decirse que la inducción, en cambio, reúne varios casos particulares que, a medida que presentan algún tipo de regularidad, pueden dar lugar a la formulación de una regla general. Por ello, es corriente escuchar que la deducción trabaja de lo general a lo particular, mientras que la inducción va de lo particular a lo particular, en busca de lo general. Sin embargo, Peirce enfatiza que, en rigor, la inducción parte ya de una teoría, que se pone a prueba; es decir, trabaja con casos particulares, pero ya con algún tipo de suposición que busca poner a prueba: "La inducción consiste en partir de una teoría, deducir de ella predicciones de fenómenos, y observar esos fenómenos para ver *qué tan cercanamente* concuerdan con la teoría" (Peirce, 1998[2012b]:282). Peirce conceptualiza la "observa-

ción" a la que hace referencia en esa cita como una investigación experimental: por ello, a diferencia de la deducción, que no tiene conexión necesaria con el mundo real, la inducción implica la conexión con el mundo empírico:

Cuando digo que por razonamiento inductivo entiendo un curso de investigación experimental, no entiendo experimento en el sentido estrecho de una operación por la que uno varía las condiciones de un fenómeno casi como uno quiera (...) Un experimento, dice Stöckhardt en su excelente *School of Chemistry*, es una pregunta que se le hace a la naturaleza. Como cualquier interrogatorio, se basa en una suposición. Si esa suposición es correcta, puede esperarse cierto resultado razonable bajo ciertas circunstancias que pueden crearse, o en todo caso, encontrarse. La pregunta es: ¿será ése el resultado? Si la Naturaleza responde "¡No!", el experimentador ha conseguido una pieza importante de conocimiento. Si la Naturaleza dice "Sí", las ideas del experimentador permanecen tal y como eran, sólo que un poco más sedimentadas. (Peirce, 1998[2012b]:281)

La persistencia en la experimentación, y la obtención de varios resultados consecutivos en los que la naturaleza diga "sí", permitirían llegar a alguna formulación de carácter más general, aunque siempre provisional, que Peirce relaciona con "el estado real de la cuestión":

La justificación para creer que una teoría experimental, que ha sido sometida a un número de pruebas experimentales, será sostenida en el futuro cercano por pruebas posteriores tales casi tan bien como lo ha sido hasta ahora es que, al proseguir firmemente en ese método, tenemos que averiguar a largo plazo el estado real de la cuestión. (Peirce, 1998[2012b]:281)

En el pensamiento de Peirce, la *abducción* aparece

como el tercer modo de razonamiento. Se trata de la adivinación, de la conjetura, de la intuición, de la elaboración de hipótesis (pero no de su comprobación). Cualquier conocimiento nuevo, dice Peirce, depende de la formación de una hipótesis, un proceso que es del orden de la abducción:

La abducción es el proceso de formar una hipótesis explicativa. Es la única operación lógica que introduce alguna idea nueva, pues la inducción no hace más que determinar un valor y la deducción meramente desenvuelve las condiciones necesarias para una hipótesis pura.

La Deducción demuestra que algo *debe ser*, la Inducción muestra que algo *es realmente* operativo y la Abducción sugiere que algo *puede ser*. (Peirce, 1998[2012b]:283).

Nos interesa particularmente detenernos en la abducción, porque parece un modo de razonamiento que podría ayudar a explicar los funcionamientos de *machine learning*, en los que los sistemas informáticos aprenden. Peirce da diversas explicaciones acerca de la facultad humana de realizar abducciones. En ciertos lugares desarrolla la idea de que la mente humana tiene una natural capacidad para imaginar correctamente algunas teorías:

Sea como sea que los hombres hayan adquirido su facultad de adivinar los modos de la Naturaleza, ciertamente no ha sido mediante una lógica auto-controlada y crítica (...) el hombre tiene una cierta iluminación (o chispazo inteligente) [*insight*] acerca de las Terceridades o elementos generales de la Naturaleza, no lo suficientemente fuerte para ser más a menudo acertada que equivocada, pero sí lo suficientemente fuerte como para no ser con una frecuencia más abrumadora equivocada en vez de acertada. (...) Esta facultad es, al mismo tiempo, de

la naturaleza general del Instinto, asemejándose a los instintos de los animales en que sobrepasa por mucho los poderes generales de nuestra razón. (Peirce, 1998[2012b]:284)

Una conceptualización de tal tipo, aunque por demás atractiva, establece una gran distancia con cualquier modelización algorítmica: ¿cómo se podrían recrear, en un sistema informático, procedimientos que consigan emular el instinto, o un “chispazo inteligente”? Esta explicación de las capacidades abductivas nos llevaría a una posición de fascinación: llegaríamos a la conclusión de que, aunque las computadoras no pueden, por definición, tener nada semejante a la intuición, el avance de la informática ha logrado, con los algoritmos de *machine learning*, emular esta facultad humana, y así hacer que los sistemas computacionales aprendan y se asemejen cada vez más a nosotros.

En otros escritos, sin embargo, Peirce explica ciertos mecanismos que operan en la elaboración de las hipótesis, que él entiende como un tipo de inferencias sintéticas. Nos interesa recuperarlos aquí, porque señalan un camino que permite conceptualizar algunos funcionamientos del *machine learning*. Dice Peirce:

26

La hipótesis se da donde encontramos alguna circunstancia muy curiosa, que se explicaría al suponer que era un caso de una cierta regla general, y en consecuencia adoptamos esa suposición, o donde encontramos que dos objetos se parecen mucho entre sí en ciertos aspectos e inferimos que se parecen mucho entre sí en otros aspectos. (Peirce, 1992[2012a]:236).

Si en la primera explicación, la abducción podía parecer, al decir de Peirce, una “conjetura ciega”, es decir, una formulación elaborada sin la intervención de ninguna observación, aquí este autor propone que

hay una suerte de “señales reveladoras”: las leemos, a veces incluso de forma no del todo consciente. Se trata de lectura de *indicios*: eso que Peirce llamaba “instinto” se apoya, en verdad, en la percepción (tal vez inconsciente) de *conexiones entre aspectos del mundo*. Estas conexiones se encuentran a partir de la identificación de semejanzas; por ello, Peirce insiste, cuando explica el funcionamiento de las hipótesis, en la cuestión de los aspectos en los que los objetos *se parecen*. La abducción sería entonces en estos casos *la capacidad de elaborar ciertas hipótesis a partir del hallazgo de semejanzas*. Como se sabe, los indicios no son, por lo general, signos seguros. Dicho de otra manera, esa conexión efectivamente observada, podría ser acertada o equivocada como explicación de la vinculación entre esos “otros aspectos” del mundo. Por ello, Peirce afirma que, si bien todo conocimiento nuevo proviene de conjeturas, estas son inútiles si no se prueban en la investigación: las hipótesis deben ser sometidas a la prueba de la inducción.

Decíamos antes que la inducción trabajaba con ya con ciertas suposiciones, que implicaban pensar alguna formulación de tipo deductivo que la experimentación inductiva pone a prueba. Se hace evidente ahora que, en verdad, esas suposiciones son hipótesis, y –por tanto– son generadas abductivamente; de allí que Peirce insista en que todo conocimiento realmente nuevo comienza necesariamente mediante una abducción. Luego, se emplea la inducción para la puesta a prueba de la hipótesis proyectada:

La abducción parte de los hechos sin, al principio, tener ninguna teoría particular a la vista, aunque está motivada por la idea de que se necesita una teoría para explicar los hechos sorprendentes. La inducción parte de una hipótesis (...) [y] necesita de los hechos para sostener la teoría. La abducción persigue una teoría. La inducción anda buscando los hechos. En la abducción la consideración de los hechos sugiere

la hipótesis. En la inducción el estudio de la hipótesis sugiere los experimentos que sacarán a la luz los verdaderos hechos a los que la hipótesis ha apuntado. (Peirce, *Collected Papers*, 7.218, citado en Sebeok y Umiker Sebeok, 1979[1994]:51)

Gérard Deladalle lo resume de este modo: “la *abducción* sugiere una hipótesis; la *deducción* extrae de ella diversas consecuencias que la *inducción* pone a prueba” (Deladalle, 1990[1996]:176). Como se ve, la abducción se ubica, en estas conceptualizaciones, en el inicio de estos procesos de conocimiento o aprendizaje, pues es a partir de ella que se generan las hipótesis que se verificarán experimentalmente en la inducción.¹³

6. MACHINE LEARNING: LÓGICAS Y FUNCIONAMIENTOS SEMIÓTICOS

Volvamos a poner ahora en foco los funcionamientos del *machine learning*, y –trabajando con el caso que hemos tomado para nuestra indagación– preguntémonos cómo podemos explicar sus funcionamientos con las lógicas descritas por Peirce.

Se puede pensar que, cuando comienzan a operar, estos algoritmos funcionan predominantemente con lógicas abductivas: inicialmente exponen los anuncios a distintos usuarios y, como vimos, algunos realizan la acción deseada, mientras que muchos no la hacen. Cuando ya hay algunos resultados positivos, el algoritmo comienza a *evaluar semejanzas* entre los usuarios. La capacidad de encontrar estos parecidos, que Peirce identificaba, entre fines del siglo XIX y comienzos del XX, como una facultad humana, hoy se encuentra también en sistemas informáticos. En estos casos, desde luego, no se explica ya por una supuesta intuición o iluminación, sino por la combinación de

un volumen masivo de datos acerca de los usuarios y sus comportamientos almacenados en forma digital (al que nos hemos referido cuando discutimos acerca de la datificación, pero que es propio del fenómeno conocido como *big data*),¹⁴ y la capacidad y velocidad de procesamiento que tienen los actuales equipos informáticos.

A partir del hallazgo de semejanzas, se puede pensar que el sistema elabora una conjetura o abducción que da lugar a una hipótesis. Como señalaba Peirce, el hallazgo de ciertas semejanzas da lugar a la pregunta acerca de si habrá semejanzas en otros aspectos. En el ejemplo: si hay cuatro personas que abrieron su cuenta en el Banco Solís, que son varones, que tienen menos de 45 años, y que poseen interés en tecnología, es probable que otros varones de 45 años o menos con interés en tecnología también abran su cuenta. Esta hipótesis se pone entonces a prueba con una reorientación de los anuncios: el algoritmo prueba qué ocurre cuando expone a otras personas de esas características a la campaña del banco. Como se ve, lo que está haciendo es ver

14 El anglicismo *big data* se emplea para referirse “a volúmenes masivos y complejos de información estructurada y no estructurada que requiere de métodos computacionales para extraer conocimiento” (Arcila Calderón, Barbosa Caro y Cabezuelo Lorenzo, 2016: 624, citado en Reviglio, María Cecilia y Castro Rojas, Sebastián, 2018:62). Como señala Ricardo Diviani (2018:16), la literatura sobre *big data* indica que uno de los empleos destacados que se realiza de estos volúmenes masivos de datos es la generación de predicciones, y en particular, la previsión de acciones futuras. En el caso que estamos viendo, los algoritmos buscan justamente predecir quiénes serán usuarios más propensos a abrir una cuenta en el banco, para decidir a quiénes le muestran la publicidad. El texto referido es Arcila Calderón, C., Barbosa Caro, E. y Cabezuelo Lorenzo, F. (2016). “Técnicas big data: análisis de textos a gran escala para la investigación científica y periodística”, en *El profesional de la información* número 25/4, páginas 623-631.

13 De hecho, Peirce señala que las buenas abducciones deberían dar lugar a hipótesis susceptibles de verificación experimental (ver Peirce, 1998[2012b]:303).

qué tan cercanamente concuerdan las respuestas de los nuevos usuarios con esa hipótesis: se trata de un funcionamiento inductivo. En este caso, es el propio algoritmo el que realiza la investigación experimental de la que habla Peirce, y es también el algoritmo el que recibe, como respuesta de los usuarios los "sí" o "no", según cliquee los anuncios y abran una cuenta en el banco, o no lo hagan. Se trata, en estos casos, de hipótesis que trabajan siempre con grados de probabilidad, buscando una maximización que nunca llega a comportarse del modo en que lo hacen las leyes con las que se maneja la lógica deductiva: no hay premisas mayores siempre verdaderas, sino maximización de probabilidades. Dicho de otra manera: no habrá ningún "segmento" para el cual la totalidad de sus usuarios descarguen la aplicación del banco y abran su cuenta; en cambio, el sistema buscará aquel segmento, o aquellos segmentos, en los cuales los usuarios con mayor probabilidad harán esta acción. Cuando alcanza la maximización de estas probabilidades, el sistema considera que ha llegado al grado máximo de aprendizaje, y predomina el funcionamiento inductivo. No obstante, el algoritmo no abandona el funcionamiento abductivo, y si, por múltiples razones, la campaña continúa y hay modificaciones en los perfiles de los usuarios que son más propensos a abrir una cuenta en el banco, volverá a experimentar inductivamente hipótesis y ajustará o calibrará el perfil de los usuarios alcanzados.

Nos referíamos recién a "segmentos" o "perfiles", lo que puede hacer pensar que la plataforma tiene, de algún modo, reunidos a los usuarios en ciertos grupos o clasificaciones. Sin embargo, lo que la plataforma tiene son datos asociados a los usuarios; ante una consulta o la acción de un algoritmo, puede "filtrarlos" y generar algún público; sin embargo, no es correcto conceptualizarlo como categorías sistemáticas, en las que los elementos (en este caso, los usuarios) se

"reparten" en conjuntos excluyentes.¹⁵ Dicho de otro modo: cada vez, el sistema hace una consulta (*query*) sobre esa base de datos de usuarios y comportamientos, que "devuelve" como resultado un conjunto de elementos que se corresponden con las condiciones determinadas. Ante cada consulta, los elementos se reordenan. En este aspecto puntual, se comportan de algún modo como otro tipo de algoritmo muy conocido: los de búsqueda por palabra clave. En ellos, todos los registros (sitios, documentos, artículos, dependiendo qué es lo que gestiona ese buscador) están "indexados", y el algoritmo devuelve un listado ordenado a partir de los términos de búsqueda ingresados por el usuario.

Los conjuntos de usuarios que, temporariamente, configuren un "segmento" o "perfil" en las campañas pueden coincidir en mayor o menor medida con una clasificación o categoría de las audiencias que tenga vida social, o con un colectivo que tenga existencia empírica en algún espacio de la vida social, más allá de su coincidencia en ciertos consumos o comportamientos digitales. En ese sentido, la pregunta por el "perfil de una audiencia" en este tipo de campañas no tiene una respuesta clara: por un lado, porque las plataformas no suelen brindar mayor información acerca de los públicos a los que les muestran los anuncios; por el otro, porque probablemente lo que llamamos la "audiencia objetivo" de la campaña sea un conjunto de registros de usuarios indexados y filtrados de acuerdo con una combinación de distintos valores asociados a determinadas variables en permanente proceso de ajuste, que sencillamente no se pueda nombrar.

7. PERSPECTIVAS

Este artículo buscó generar una aproximación a la

¹⁵ En otro trabajo, hemos señalado que la tensión entre estos distintos modos de gestionar los públicos receptores en Internet se manifiesta en el par audiencias/usuarios (ver Zelcer, 2014).

conceptualización de los procesos de aprendizaje automático por computadora desde una perspectiva semiótica peirciana. En ese sentido, el territorio de la publicidad digital en general, y el caso de las campañas del Banco Solís en particular, sirvieron como excusa para trabajar en ese territorio mayor. A modo de resumen, a lo largo del trabajo hemos descripto cómo, trabajando con *big data*, los algoritmos de *machine learning* modelan y buscan recrear esa “facultad humana” señalada por Peirce de la adivinación o la conjetura conocida como abducción, e identificamos que el hallazgo de semejanzas por parte del sistema es una de las operaciones básicas sobre las que se apalanca el aprendizaje automático, que se despliega con procedimientos que combinan lógicas abductivas con otras inductivas, y cuyos resultados realimentan al mismo funcionamiento del algoritmo.

Estas observaciones acerca del *machine learning* exceden el mundo de la publicidad digital: funcionamientos similares se observan en otras plataformas y sistemas informáticos que gestionan datos de usuarios. Y al mismo tiempo, operan también en sistemas que gestionan –retomando a Gendler– datos ya no únicamente de sujetos, sino de sujetos, objetos y prácticas, como las principales plataformas de consumo musical y audiovisual que existen en la actualidad. En ese sentido, este artículo busca abrir un camino de indagación que pueda continuar desarrollándose y aplicándose a otros fenómenos contemporáneos de los consumos culturales digitales.

BIBLIOGRAFÍA

- Bazzara, Lucas (2021). “Datificación y streamificación de la cultura. Nubes, redes y algoritmos en el uso de las plataformas digitales”, en *InMediaciones de la Comunicación*, 16(2), 37-61. Disponible en <https://revistas.ort.edu.uy/inmediaciones-de-la-comunicacion/article/view/3082> [Consulta: 9/2/2022]
- Bonini Tiziano y Gandini, Alessandro (2019) “First week is editorial, second week is algorithmic: platform gatekeepers and the platformisation of music curation”, publicado en *Social Media + Society* número 5(4). Disponible en https://www.researchgate.net/publication/337432337_First_Week_Is_Editorial_Second_Week_Is_Algorithmic_Platform_Gatekeepers_and_the_Platformization_of_Music_Curation [Consulta: 19/2/2022].
- Cheung, Steven, et. al (2018) “A functional perspective on machine learning via programmable induction and abduction”. Disponible en https://www.researchgate.net/publication/323856247_A_Functional_Perspective_on_Machine_Learning_via_Programmable_Induction_and_Abduction [Consulta: 29/9/2019].
- Cingolani, Gastón (2017a) “Sistemas de recomendación, mediatizaciones de lo preferible y enunciación”, en Busso, Mariana Patricia y Camuso, Mariángeles (eds.) (2017) *Mediatizaciones en tensión: el atravesamiento de lo público*, Rosario, UNR Editora. Páginas 30 a 47
- _____ (2017b) “Estrategias para el acceso: los sitios de recomendación como espacios de tensiones en la circulación y mediatización del reconocimiento”, en *A circulação discursiva: entre produção e reconhecimento*, Paulo César Castro (org.), Maceió, EDUFAL, pp. 125-140. ISBN 978-85-5913-125-3
- Deladalle, Gérard (1996) *Leer a Peirce hoy*, Barcelona, Gedisa. Edición original (1990) *Lire Peirce Aujourd'hui*, Bruselas, Boeck-Wesmael.
- Diviani, Ricardo (2018) “Consideraciones epistemológicas, teóricas y críticas en relación al big data” en *La mediatización contemporánea y el desafío del big data* en Biselli, Rubén y Maestri, Mariana (editores) (2018) *La mediatización contemporánea y el desafío del big data*, Rosario, UNR Editora.
- Hurwitz, Judith y Kirsch, Daniel (2018) *Machine Learning For Dummies®*, IBM Limited Edition, New Jersey, John Wiley & Sons, Inc.

- Koldobsky, Daniela (2019) "Búsqueda de música en plataformas mediáticas: entre el archivo y la novedad". Ponencia presentada en el 4º Congreso de la Asociación Internacional de Semiótica, Buenos Aires, 13 de septiembre de 2019
- Maluenda de Vega, Raquel (2021) "Qué es un algoritmo informático: características, tipos y ejemplos", en <https://profile.es/blog/que-es-un-algoritmo-informatico/> [Consulta: 24/1/2022]
- Mehler, Alexander (2003) "Methodological aspects of computational semiotics", S.E.E.D. Journal (Semiotics, Evolution, Energy and Development), 3(3), páginas 71-80.
- Peirce, Charles Sanders (1986) *La ciencia de la semiótica*, Buenos Aires, Nueva Visión. Edición original: Hartshorne, Charles y Weiss, Paul (comp.) (1965) *Collected Papers of Charles Sanders Peirce*, Cambridge-Massachusetts, Harvard University Press. Volumen II, *Elements of Logic*, libro II, "Speculative Grammar", caps. 1, 2 y 3.
- _____ (2012a) *Obra filosófica reunida. Tomo I (1867-1893)*, México D.F., Fondo de Cultura Económica. Traducción de Darin McNabb. Edición original: (1992) *The Essential Peirce. Selected Philosophical Writings. Volume 1 (1867-1893)*, Indiana, Indiana University Press.
- _____ (2012b) *Obra filosófica reunida. Tomo II (1893-1913)*, México D.F., Fondo de Cultura Económica. Traducción de Darin McNabb. Edición original: (1998) *The Essential Peirce. Selected Philosophical Writings. Volume 2 (1893-1913)*, Indiana, Indiana University Press.
- Reviglio, María Cecilia y Castro Rojas, Sebastián (2018) "El cuerpo del corpus #RosarioSangra. Colaboración artesanal-computacional para el estudio de la articulación de regímenes de visibilidad" en Biselli, Rubén y Maestri, Mariana (editores) (2018) *La mediatización contemporánea y el desafío del big data*, Rosario, UNR Editora.
- Ricci, Francesco, Rokach, Lior y Shapira, Bracha "Recommender Systems: Introduction and Challenges", en Ricci Francesco, Rokach Lior, Shapira Bracha y Kantor, Paul B. (eds.) *Recommender Systems Handbook*, 2nd edition, Springer: New York, 2015.
- Scolari, Carlos (2018) *Las leyes de la interfaz*, Barcelona, Gedisa.
- Sebeok, Thomas A. y Umiker-Sebeok, Jean (1994) *Sherlock Holmes y Charles S. Peirce. El método de la investigación*, Barcelona, Paidós. Edición original: (1979) *You Know My Method*, Indiana, Gaslit Publications. Traducción de Lourdes Güell.
- Wilmott, Paul (2019) *Machine Learning. An applied mathematics introduction*, Panda Ohana Publishing
- Zelcer, Mariano (2014) "Audiencias/usuarios: una discusión inicial acerca de las categorías empleadas para hablar sobre la recepción en Internet.", en L.I.S. Letra. Imagen. Sonido. Ciudad mediatizada, año VI, número 12, 15-28. Buenos Aires: UBACyT. Ciencias de la Comunicación, Facultad de Ciencias Sociales, UBA. Disponible en <https://publicaciones.sociales.uba.ar/index.php/lis/article/view/3773> [Consulta: 19/2/2022]
- _____ (2019) "Publicidad en la Web: de la lógica de los medios masivos a los anuncios personalizados", publicado en *Designis* número 30 (enero-junio 2019). Disponible en https://ddd.uab.cat/pub/designis/designis_a2019n30/designis_a2019n30p123.pdf [Consulta: 19/2/2022]

DATOS DEL AUTOR

Mariano Zelcer

Argentina

IIEAC (Instituto de Investigación y Experimentación en Arte y Crítica), Área Transdepartamental de Crítica de Artes Oscar Traversa, Universidad Nacional de las Artes (UNA)

E-mail: marianozelcer@yahoo.com.ar

Áreas de investigación o interés: Semiótica, dispositivos, imágenes digitales, machine learning.

Filiación Institucional: IIEAC (Instituto de Investigación y Experimentación en Arte y Crítica), Área Transdepartamental de Crítica de Artes Oscar Traversa, Universidad Nacional de las Artes (UNA)

Dirección postal: Lambaré 959, 7º 20 (1185), Ciudad de Buenos Aires, Argentina.

Teléfono: (011) 15-5584-8202

Fecha: 20/2/2022

ORCID: 0000-0001-7472-0906

REGISTRO BIBLIOGRÁFICO

Mariano Zelcer. "Machine learning y lógicas semióticas: el caso de la publicidad digital" en *La Trama de la Comunicación*, Vol. 26 Número 2, Anuario del Departamento de Ciencias de la Comunicación. Facultad de Ciencia Política y Relaciones Internacionales, Universidad Nacional de Rosario. Rosario, Argentina. UNR Editora, enero a junio de 2022 p. 015-031. ISSN 1668-5628 – ISSN 2314-2634 (en línea).

RECIBIDO: 20/02/2022

ACEPTADO: 27/02/2022