

5 step per capire se si sta leggendo una fake news su Facebook



Sui social network è facile imbattersi nelle fake news. Ecco alcuni consigli semplici e immediati per provare a smascherare una bufala

Sentiamo parlare di fake news quotidianamente e ancora più spesso le ritroviamo nel nostro feed di Facebook o condivise con un tweet. Ormai hanno invaso i social network e la loro diffusione sta andando fuori controllo. Come può difendersi un utente di internet? Basta poco, in fondo. Bisogna imparare a **condividere la notizie consapevolmente** e mettere quindi in atto quello che in gergo si chiama *fact checking*, ovvero il lavoro di verifica che ogni buon giornalista deve fare per accertarsi che gli avvenimenti citati o i dati usati in un determinato articolo siano veri.

Fare un controllo prima di una condivisione non è affatto

complicato. **Ecco una guida in 5 passaggi.**

1. Appena leggi una notizia che ti sembra clamorosa confrontala. Il modo più immediato è **controllare su Google o su altre fonti online** se è stata ripresa e se sia veritiera. In alcuni casi vedrai che alcuni siti che si occupano di smascherare bufale ne hanno già parlato e **l'hanno già smentita.**

2. Se se nessuno ha smentito un fatto, non significa comunque che sia vero.

Guardate in modo scrupoloso il post per **capire se c'è qualche anomalia.** Se si tratta di un post con un link a un sito, **controllate la testata:** esistono infatti molti siti di fake news che hanno nomi molto somiglianti a quelli di testate giornalistiche note. Sono creati ad hoc per ingannare la gente che legge distrattamente.

3. Arrivati al punto 3 è ora di metterci un po' più di tempo e di impegno. Capita spesso infatti di leggere, soprattutto su Facebook, solo notizie con una breve descrizione sotto una fotografia, condivise direttamente da una pagina del social network e senza link esterni.

In questi casi può essere d'aiuto **Google con la sua ricerca per immagini.** Salvate l'immagine dal contenuto sospetto, andate sul motore di ricerca, cliccate l'icona della macchina fotografica e poi su "*carica un'immagine*", quindi inserite la fotografia. Quasi certamente vi apparirà l'immagine con la reale descrizione.

4. **Attenzione alla data della notizia.** Spesso a una notizia vera vengono affiancate immagini altrettanto vere ma che non si riferiscono a una news. In questo caso potete sempre provare a inserire l'immagine su Google come spiegato nel passaggio precedente o se preferite potete utilizzare [Tine Eye](#), un sito che permette non solo di scoprire di che immagine si tratta ma anche quando e su quale sito è stata utilizzata online.

5. Se avete controllato tutto quanto sopra allora provate a verificare se davvero la fonte è attendibile e cliccate sul link della notizia, nel caso si tratti di un *page post link*, o sulla pagina che ha condiviso il contenuto, nel caso si tratti di un testo con immagine. Nel primo caso **guardate bene in che tipologia di sito vi trovate**, cercate se ci sono i credit del sito web, se si tratta di una testata giornalistica e quali

sono le altre notizie che sono condivise. Se sono tutte sensazionalistiche e non vedete credit in chiaro potreste essere di fronte a un **sito di bufale**. Lo stesso vale per la pagina Facebook.

Ecco cosa fare per evitare il diffondersi di bufale e notizie false, che alimentano sciacallaggi mediatici e, ancora peggio, fanno guadagnare chi crea notizie false a tavolino.

Come avrete notato, vi consigliamo di cliccare sulla notizia soltanto arrivati al quinto step. Perché? Per **evitare di far guadagnare chi lucra** sulle fake news. Questi siti guadagnano grazie al grande numero di visualizzazioni che riescono a ottenere. In media il guadagno che possono ottenere è di 2 euro ogni mille visualizzazioni e, grazie anche ad una condivisione spensierata, una bufala può essere letta anche da 500.000 persone, portando un guadagno di 1.000/1.500 € a chi le ha pubblicate.

Ogni volta che condividiamo una notizia sulla nostra **bacheca di Facebook** o su **Twitter**, contribuiamo da una parte alla diffusione di informazione sbagliata e dall'altra ad aumentare i guadagni di chi questa disinformazione la cavalca. Se vi è capitato erroneamente di condividere una fake news rimuovete il contenuto così da impedire che altre persone leggendo la facciano come voi e clicchino "*share*" alla leggera, potrà sembrare un'azione da poco ma la grande visibilità che queste notizie ricevono è formata proprio da tanti **piccoli share** inconsapevoli.

Gli algoritmi e i "mostri" della tecnologia



Ho visto lo sventurato, il miserabile mostro che avevo creato (...) demoniaco cadavere a cui avevo così miseramente dato la vita (Mary Wollstonecraft Godwin)

Nel 2007 l'allora sindaco di Washington, Adrian Fenty, era determinato a risolvere un problema: gli studenti di molte scuole non raggiungevano buoni risultati. Si partì dall'idea che la colpa doveva essere degli insegnanti, che non facevano bene il loro lavoro. Per cui si affidò il compito di valutare gli insegnanti ad un software denominato IMPACT.

L'algoritmo prendeva in considerazione una serie di elementi, tra i quali dei test da somministrare agli studenti e da ripetere nel tempo. Il modello, detto a valore aggiunto (value added model) è fondamentalmente un modello statistico che cerca di distinguere l'impatto causale di un insegnante sull'apprendimento dei suoi studenti da altri elementi, quali abilità degli studenti e fattori extrascolastici.

Gli algoritmi si espandono

I modelli statistici per la valutazione degli insegnanti negli Usa risalgono agli anni '70. Man mano che la potenza di calcolo degli elaboratori aumentava, questi sistemi

automatizzati si sono diffusi e perfezionati. Ma il primo settore nel quale sono stati applicati gli algoritmi è probabilmente quello finanziario. I sistemi automatizzati hanno rivoluzionato l'intero settore, finendo per sostituire gli esseri umani nelle "decisioni" su vendite e acquisti. Poi, nel 2008, il crollo della borsa fece comprendere che non solo gli algoritmi erano sempre più connessi al mondo degli uomini, ma che potevano anche alimentare ed amplificare i problemi della società.

Negli anni seguenti, come se nulla fosse accaduto, sorsero algoritmi sempre più sofisticati, intelligenze artificiali che imparavano da sole (machine learning) e prendevano decisioni in totale autonomia. Il loro uso si è progressivamente espanso a tutti i settori della società. La statistica e la matematica iniziarono ad occuparsi attivamente di comportamenti, desideri, movimenti, acquisti, crimine e terrorismo.

La presenza di algoritmi nelle nostre vite ormai è pervasiva. Un algoritmo può essere usato (Nicholas Diakopoulos, Algorithmic accountability reporting: on the investigation of black boxes) a fini di prioritizzazione, per stabilire come devono essere distribuite le risorse dei servizi. Ad esempio, per le ispezioni agli edifici per verificare se sono adeguati i servizi antincendio, oppure se reggono in caso di terremoto. La polizia americana utilizza spesso algoritmi di prioritizzazione (es. Predpol) per stabilire in quali zone della città deve inviare le pattuglie disponibili.

Un altro utilizzo comune degli algoritmi è a fini di classificazione, ad esempio per distinguere tra contenuti leciti o illeciti online. Un algoritmo di associazione, invece, crea relazione tra elementi, come se fossero hyperlink. Serve a raggruppare elementi, ad esempio i pregiudicati che hanno avuto problemi con droghe. Infine abbiamo gli algoritmi di filtraggio, che includono o escludono informazioni in base a specifiche regole o criteri. Ovviamente ogni algoritmo può avere anche più scopi.

Un algoritmo di filtraggio può, quindi, essere anche un algoritmo di classificazione. Come ContentID di Youtube che

classifica i contenuti presunti illeciti online e li filtra. Oppure come una App di news personalizzate che classifica le notizie, le filtra in base a criteri prefissati dagli utenti (utilizzando i criteri inseriti dal programmatore) e quindi le associa agli utenti, prioritizzandole.

Il passo dalla valutazione alla predizione è stato relativamente breve. Oggi gli algoritmi predicono la nostra attendibilità, il nostro potenziale, il nostro futuro. C'è un algoritmo per stabilire se siamo (saremo?) bravi insegnanti, un altro per prevedere se saremo buoni studenti, e poi lavoratori, mariti o mogli, amanti, criminali, perfino terroristi.

Sarah Wysocki, insegnante della MCFarland Middle School, era unanimemente considerata tra le migliori. La descrivevano entusiasta, creativa, visionaria, flessibile, motivante e incoraggiante. Eppure venne licenziata, insieme ad oltre 200 insegnanti, a causa di un pessimo punteggio calcolato da Impact. I bassi punteggi ottenuti dai suoi studenti delle elementari forse erano dovuti alle problematiche del quartiere popolare, forse alla povertà delle famiglie, forse anche a problemi di salute dei ragazzi. Ma per i programmatori di Impact non erano motivi validi. La fine della sua carriera è stata decisa da un algoritmo.

Black Box

Un algoritmo consta di tre elementi: i dati di input, il processo algoritmo vero e proprio, il risultato, cioè l'output, si tratti di una previsione o un punteggio. Di questi tre elementi molto spesso si conosce solo l'output. L'algoritmo vero e proprio (il codice) non è conoscibile perché è una "proprietà intellettuale", ed è protetta anche in base alla recente direttiva Trade Secrets dell'Unione europea. Le nostre vite sono sempre più esposte e trasparenti, in base a leggi introdotte dai nostri governi e logiche di mercato. Il regime di "sorveglianza" al quale siamo sottoposti pervasivamente è giustificato da ragioni di "sicurezza", o

semplicemente per motivi di “efficienza economica” o per “semplificare la navigazione online”.

Nel contempo i metodi operativi delle aziende e delle stesse istituzioni sono sempre più opachi, gli algoritmi non sono trasparenti ma vere e proprie black box (Frank Pasquale, *The black box society*). La segretezza degli algoritmi è tale da impedirci di comprenderne le logiche e quindi di distinguerne il buon funzionamento dall'abuso. Eppure, sono gli algoritmi che ormai dettano le regole di tantissime attività umane, come ad esempio nella selezione del personale.

Il problema è che gli algoritmi, prima di prendere decisioni autonomamente, devono essere addestrati. Il training prevede che siano forniti al sistema degli input, opportunamente identificati (taggati), in modo che l'algoritmo impari cosa riconoscere.

Nel 2009 (*HP Computers are racist*) un software di rilevamento facciale non riesce a riconoscere una persona di colore. Secondo l'azienda non era un problema “razziale”, ma una questione di insufficiente illuminazione dello sfondo che determinava problemi di contrasto. Il software aveva difficoltà a “vedere” persone di pelle scura. In realtà un software non “vede”, ma umanizzandolo l'azienda nascondeva il vero problema, e cioè che il software di fatto determinava, ovviamente non intenzionalmente, uno svantaggio per le persone di pelle scura. In sostanza non era “neutrale”. Quello che per l'azienda era un problema meramente “tecnico”, poteva però avere delle implicazioni nella vita reale sotto forma di discriminazione sociale. La causa probabilmente stava nel fatto che nel laboratorio la maggior parte dei dipendenti erano bianchi.

Altro caso interessante riguarda un sistema algoritmico che produceva risultati “sessisti”, associando alle donne immagini di cucina e così via. Analizzando gli input forniti alla macchina, si vide che due collezioni di immagini, tra cui una supportata da Microsoft e Facebook, presentavano una distorsione di genere nella raffigurazione di attività come la cucina e lo sport: mentre le immagini di cucina, shopping e

lavaggio erano associate a donne, quelle di sport erano legate ad uomini. Il software di apprendimento automatico in fondo non faceva altro che il suo lavoro: apprendeva.

L'esperimento più interessante è quello del chatbot Tay della Microsoft. Doveva essere un esperimento intelligente di apprendimento automatico, imparando direttamente dalle persone con le quali interagiva su Twitter. L'idea era che diventasse indistinguibile dai millennial, ma ci sono volute meno di 24 ore perché Tay si trasformasse in un razzista xenofobo. Secondo alcuni, in realtà, l'esperimento è stato un successo. Il bot, preda dei troll, ha imparato fin troppo bene dagli utenti. La debacle è l'esempio di come gli esseri umani possono corrompere la tecnologia.

Gli algoritmi sono costruiti per approssimare il mondo in modo da soddisfare gli scopi del loro architetto e incorporano una serie di presupposti su come funziona e su come dovrebbe funzionare il mondo. Gli algoritmi, quindi, possono riflettere i pregiudizi dei programmatori incorporati nel codice, nel momento in cui interpretano e leggono i dati. Perché un programmatore generalmente non riconosce come pregiudizievoli i suoi criteri e quindi può inavvertitamente (ma anche volontariamente, volendo) introdurre dei criteri non neutrali. Il processo decisionale dipendente da algoritmi finirà per replicare i pregiudizi strutturali su vasta scala, ampliandoli.

Un software di screening valuta i candidati in base a criteri che vengono inseriti nel codice, criteri spesso soggettivi, come il nome. Se un nome "suona" da bianco piuttosto che da nero, da italiano piuttosto che straniero, da settentrionale piuttosto che meridionale, da uomo piuttosto che donna (Amazon e l'intelligenza artificiale sessista: non assumeva donne). Perché un algoritmo sia davvero "neutrale" occorre che i dati di training siano neutrali anch'essi, che i criteri di selezione inseriti nel codice siano neutrali.

Le forme di discriminazione ipotizzabili sono molte. Alcune polizze assicurative potrebbero essere limitate in base al luogo in cui si vive, le migliori condizioni per le carte di

credito potrebbero essere offerte solo a determinate persone, un negozio online potrebbe mostrare prezzi diversi per gli stessi prodotti in base alle caratteristiche individuali, un prestito potrebbe essere rifiutato a chi ha amici poveri sui social, l'applicazione di misure di prevenzione potrebbe avvenire in base ai vicini di casa.

I bias degli algoritmi

Sui bias (pregiudizi) degli algoritmi, Carlo Blengino fa una interessante riflessione su Il Post, che conclude con una serie di domande:

Può una macchina, un algoritmo evoluto, aiutare, correggere e fin anche sostituire il giudizio umano? Ma soprattutto, i nostri inevitabili bias cognitivi saranno amplificati o troveranno invece mitigazione, se non soluzione, nell'uso di sistemi esperti artificiali?.

Gli esseri umani, è noto, sono intrisi di pregiudizi, è il meccanismo di funzionamento del cervello che alimenta tali bias. Ciò che differenzia un essere umano da una scimmia è il meccanismo di compressione cognitiva col quale analizziamo le informazioni in blocchi piuttosto che singolarmente, che è alla base della creazione delle procedure complesse e automatizzate che ci governano (abitudini). In tal modo il cervello può concentrarsi su più cose contemporaneamente rendendoci molto più efficienti. Di contro, molte informazioni non sono più analizzate dalla coscienza critica.

La capacità di riconoscimento dei pattern e la capacità combinatoria, ci hanno permesso di essere la specie animale più avanzata, ma porta anche un lato oscuro. Quel processo comporta la possibilità di saltare a conclusioni sbagliate per il semplice motivo che le informazioni iniziali non vengono sempre sottoposte ad analisi. Siamo così desiderosi di cercare dei modelli nel mondo che ci circonda, e così soddisfatti quando li abbiamo trovati, che non effettuiamo controlli

sufficienti sulle nostre intuizioni apparenti.

E gli inevitabili bias degli essere umani possono essere introdotti negli algoritmi nella fase di programmazione o di addestramento. Con l'amplificazione dovuta alla scala, enormemente più vasta, di applicazione dell'algoritmo.

Quindi, un algoritmo aiuterà l'uomo a liberarsi dei suoi bias? Oppure sarà l'uomo a risolvere i bias degli algoritmi con maggiori controlli sulla neutralità dei codici?

Una “soffiata” digitale

Una fonte confidenziale chiama la Polizia, comunicando che in un certo luogo vi sono due persone intente a spacciare droghe. La Polizia si reca sul posto. Dopo solo due minuti di osservazione, un evento imprevisto (sopraggiunge qualcuno che potrebbe vederli?) decidono di intervenire. Hanno visto due persone, sono lì ferme. Nonostante vi fossero degli elementi per ritenere che una delle due persone fosse solo un acquirente, entrambe vengono arrestate per spaccio.

La “fonte confidenziale”, infatti, parlava di “due” spacciatori. Giunti sul posto, i poliziotti vedono due persone, e, poiché la fonte confidenziale era stata “riscontrata” (cioè trovata attendibile) più volte, “leggono” gli elementi che si presentano ai loro occhi in base alla “soffiata”.

Secondo il “principio di coerenza”, infatti, si tende a scegliere o a creare l'interpretazione dei fatti più coerente con i dati disponibili. Il principio del “minimo impegno” fa sì che il soggetto tenda a compiere solo le inferenze minime indispensabili a produrre un'interpretazione coerente, per cui tra due spiegazioni ugualmente coerenti sceglierebbe quella che richiede il minor numero di supposizioni. Elementi tali da far ritenere che una delle due persone è solo un acquirente non vengono considerati, oppure sono ridotti a coerenza con le informazioni fornite dalla “fonte confidenziale”.

Autorità e controllo

Il processo di interpretazione dei dati è un processo complesso che può portare ad errori, specialmente in casi dubbi, nei quali i dati possono essere interpretati in modi diversi. Il modo di rapportarsi dell'“interprete” ai dati è fondamentale, perché un pregiudizio dell'interprete può portare a risultati sbagliati.

Se il risultato, cioè l'interpretazione dei dati, viene da un algoritmo, l'output potrebbe essere discriminatorio, se l'algoritmo non è neutrale. Ma il problema primario è che per verificare se l'algoritmo è neutrale o meno è innanzitutto necessario essere coscienti del fatto che l'algoritmo possa sbagliare. Ma, chi utilizza i risultati dell'algoritmo spesso non è a conoscenza del codice di programmazione né dei dati di input (addestramento). Se l'algoritmo è stato utilizzato varie volte e in molti casi è risultato attendibile (è stato “riscontrato”) si tende a credere che lo sia sempre. Il poliziotto che decide di inviare una pattuglia in una determinata zona in base all'algoritmo difficilmente penserà che l'algoritmo si può sbagliare.

Una black box in fondo, non è altro che una soffiata digitale. Se riscontrata più volte la si ritiene affidabile. Diventa incontestabile, una sorta di “autorità” che non è possibile mettere in discussione.

Come ha evidenziato Milgram con i suoi esperimenti, le persone sono estremamente disponibili ad assoggettarsi a un'autorità. E l'intera società è strutturata in modo da condizionare, attraverso le istituzioni (scuola, famiglia, lavoro), l'individuo a sottostare alla regola dell'obbedienza all'autorità. Ma oggi l'autorità è sempre più espressa in forma di algoritmo.

È l'impenetrabile black box, difesa dalle leggi moderne, a garantire l'esercizio del potere nelle società moderne. La strutturazione del sistema in black box rimuove le responsabilità dei singoli rendendo difficile nell'ambito delle gerarchie e delle competenze stabilire chi deve pagare

per un eventuale errore. L'obbedienza all'autorità fa il resto.

Sulla segretezza degli algoritmi si fonda l'elemento essenziale della "sorveglianza", cioè il "controllo". L'individuo viene frammentato in dati e ricomposto a creare quella che non è la banale transcodifica della sua vita in informazioni gestibili da un elaboratore, quanto piuttosto l'esercizio di un vero e proprio potere, una forma di controllo (John Cheney-Lippold, We are data). Essere governati non è altro che essere controllati costantemente, osservati, indottrinati, valutati e censurati. Se un medico per svolgere il suo lavoro necessita di un badge che gli consente l'ingresso nel reparto, vuol dire che c'è un algoritmo che stabilisce quando e se quel medico può essere un medico. Basta revocargli i permessi perché non lo sia più. Così per uno studente, e per un insegnante. È un algoritmo a decidere chi o cosa siamo. Dipende dall'algoritmo se noi abbiamo dei diritti, perché in fondo dipende dall'algoritmo se siamo identificati come soggetti che hanno diritti. Questa è la "datificazione". Oggi una persona è datificata non tanto in base a ciò che è, ma per come è vista dal programmatore. L'esempio più classico: il "terrorista". Un terrorista oggi non è più una persona, quanto piuttosto un modello (type) ricavato dalla datificazione di una serie di comportamenti tipici di soggetti ritenuti terroristi (signature, o firma). I dati (comportamenti) da estrarre sono selezionati dai programmatori del modello, quindi alla fine il modello di terrorista non descrive affatto un terrorista quanto piuttosto come un terrorista è visto dal programmatore (lo schema di un terrorista).

Ad esempio, se delle persone sono viste con armi in prossimità di un luogo frequentato da terroristi, oppure se sono viste comunicare con terroristi, o dare ordini a terroristi. È anche possibile che l'algoritmo confonda e riunisca in un unico individuo il comportamento di più soggetti che abitano lo stesso appartamento. La costruzione di un modello ideale nel quale far rientrare una persona fisica non si basa sulla

realtà quanto piuttosto è un'approssimazione di un fenomeno dinamico, la traduzione in una quantità di numeri trattabili da un software. In tal senso la costruzione di un modello di terrorista (ma anche di altre tipologie) non è tanto l'estrazione di dati dalla realtà (raw data), quanto piuttosto la costruzione di dati a partire dall'osservazione della realtà (cooked data). E, come tale, è soggetta a molteplici errori e fenomeni discriminatori.

La classificazione in base ad algoritmi non riscontra l'individuo effettivo come si estrinseca nella vita reale, quanto piuttosto un suo "doppio" virtuale, che può essere "maschio" ma anche "femmina", "famoso" o "non famoso", "giovane" o "vecchio". La categorizzazione di "uomo" non è un uomo bensì la datificazione dello schema di attività che generalmente pone in essere un "uomo". Per un algoritmo una donna che si "comporta" da uomo, legge riviste da "uomo", fa acquisti da "uomo", è un "uomo", con un disallineamento tra la vita reale e la vita basata sui dati (vita online o virtuale). E la classificazione è del tutto indipendente dalla volontà dell'individuo. Nella sua visione più ottimistica, quella "Silicon Valley oriented", i problemi degli algoritmi sono facilmente risolvibili. Come? Accumulando più dati. Sempre più dati.

Nella sua visione più ottimistica, quella "Silicon Valley oriented", i problemi degli algoritmi sono facilmente risolvibili. Come? Accumulando più dati. Sempre più dati.

Cosa è normale?

I modelli algoritmici utilizzano i dati dei comportamenti passati per prevedere il futuro. In tal modo si realizza una serie di schemi comportamentali (profili) che servono a classificare le persone fisiche. I cittadini vengono categorizzati assegnando loro un profilo di rischio. L'inclusione in categorie dipende e rafforza l'idea che certi comportamenti siano la "norma".

Porre in essere comportamenti che si discostano da tale

“normalità”, comporta automaticamente l’inclusione in categorie a rischio, finendo per indicare anche l’avvio di un percorso delinquenziale o terroristico. E il discrimine tra attività sovversiva e semplice attività di protesta è molto sottile agli occhi dello stato.

Così, il sistema finisce per discriminare tutti coloro che non si comportano a “norma”, favorendo il conformismo. Tende a replicare possibili disuguaglianze (si pensi ai musulmani, ritenuti spesso non a norma in paesi occidentali), oltre a creare difficoltà nell’affrontare eventi del tutto nuovi od inaspettati.

Monstrum

“Ho visto lo sventurato, il miserabile mostro che avevo creato (...) demoniaco cadavere a cui avevo così miseramente dato la vita”.

Così Mary Wollstonecraft Godwin introduce la “creatura” nel libro “Frankenstein o del Prometeo moderno”. Il suo artefice crea il mostro perché lo riverisca come un dio. Ma il mostro (dal latino “monere”, avvertire) è anche una critica spietata, non tanto alla scienza e alla tecnologia, ma a ciò che con essa viene fornita (Anne k. Mellor, Mary Shelley: Her Life, Her Fiction, Her Monsters).

Il mostro, in realtà, è un’inevitabilità storica, una necessità per stabilire i confini tra la società e ciò che la minaccia, ciò che deve rimanerne fuori, ciò che le è alieno. Nel romanzo al di fuori dei confini della città (della società) stanno i demoni, fatti di parti di cadavere, di frammenti e pezzi presi qua e là e ricomposti in un tutt’uno (come i dati raccolti e ricomposti dagli algoritmi), per creare ciò su cui devono essere focalizzate le paure e i timori della società intera (così come oggi si fa con gli islamici, i migranti, i terroristi, ecc...). Il mostro rappresenta i mali della società, e una volta cacciato, ci assicurano, i problemi saranno tutti risolti. Ma la scomoda

verità è che il mostro non è affatto a noi estraneo o diverso o separato, bensì è la tecnologia (la creazione) a separarlo dal resto della società. Al punto che senza di essa il mostro non ha più motivo di esistere (nel romanzo senza il suo creatore il mostro si uccide).

“Cosa è normale in un mondo retto dagli algoritmi?”, si chiede John Cheney-Lippold in *We are data*. Nella società capitalista, in quella moderna retta da algoritmi, il mostro è tutto ciò che non rientra nella normalità, che non rientra nelle categorie accettate nella società, tutto ciò che è diverso, coloro che non si assoggettano all'autorità costituita, quelli che sfuggono alla black box. Creare il mostro ci consente di definire i confini di cosa è normale e cosa no. Nella società moderna, dove regna l'imperativo tecnologico, mirabilmente riassunto nelle parole del nipote del professore in *Frankenstein Junior* di Mel Brooks (“si può fare!”), per cui tutto ciò che è potenzialmente fattibile va fatto senza tenere conto delle implicazioni etiche, il compito di creare il “mostro” è svolto sempre più dagli algoritmi.

E Sarah Wysocky, l'insegnante licenziata a causa di un pessimo punteggio calcolato da un algoritmo? L'algoritmo che la licenziò fu contestato. Un modello statistico dovrebbe essere basato su grandi quantità di dati, ma nel caso specifico il tutto si riduceva ai test di una trentina di studenti. In particolare si evidenziò che il modello non funzionava correttamente perché la distribuzione degli studenti non era casuale. Piuttosto dipendeva da fattori quali le richieste specifiche dei genitori, dalle specializzazioni degli insegnanti, e così via. Un insegnante ritenuto particolarmente portato per bambini con problemi di concentrazione otteneva inevitabilmente punteggi più bassi. Così la selezione in base all'algoritmo finiva per creare forti disincentivi per gli insegnanti ad occuparsi degli studenti più problematici.

Ciò non portò ad alcun ripensamento del sistema, che anzi fu giustificato anche di fronte a evidenze contrarie. Buoni punteggi ottenuti dagli studenti di insegnanti con basso punteggio erano ritenuti confermativi della valutazione del

sistema perché in tali casi l'insegnante, a rischio di licenziamento, era portato ad alterare i punteggi degli studenti. Infine, si concluse, il modello statistico può anche sbagliare in casi singoli, ma l'efficacia si vede nel complesso. In sostanza non era il sistema che si adattava alla realtà, ma in alcuni casi si adattava la realtà per giustificare il sistema.

Sarah Wysocki fu licenziata. Ma, essendo comunque un'ottima insegnante venne immediatamente assunta da un'altra scuola. In un quartiere per ricchi, che non utilizzava quel modello di valutazione degli insegnanti.

Brasile, allarme fake news su WhatsApp alle presidenziali. Il Tribunale elettorale: «Ci vogliono misure severe»



Secondo il sito anti-bufale Aos Fatos soltanto le principali dodici notizie false circolate nel fine settimana delle elezioni sono state condivise 1,17 milioni di volte su Facebook

RIO DE JANEIRO – L'influenza di WhatsApp sulla corsa presidenziale, con il trionfo delle fake news, ha raggiunto un livello non più tollerabile e andrebbe contrastata, sostiene il Tribunale elettorale brasiliano. Bisogna «stabilire una qualche forma di controllo sul flusso di informazioni per impedire che si ripeta quanto accaduto al primo turno, quando bugie e montature hanno causato danni notevoli ad alcuni candidati», scrive il consiglio di esperti dell'Authority elettorale. Con alcuni dei suoi membri suggerendo addirittura «misure dure». Inutile continuare a ignorare la realtà, dicono: WhatsApp da «messaggero» di informazioni e opinioni personali si è trasformato esso stesso in un social network, e deve essere responsabilizzato per i propri contenuti. L'opinione è però assai controversa.

Elezione falsata?

Secondo il sito [Aos Fatos](#), implacabile cacciatore di bufale in Rete, soltanto le principali dodici notizie false circolate nel fine settimana delle elezioni sono state condivise 1,17 milioni di volte su Facebook. Campione di clic un video taroccato dove si vede l'urna elettronica con la quale si vota in Brasile che da sola compone il numero del candidato di sinistra Fernando Haddad ([al ballottaggio del 28 ottobre col 28% dei voti](#)), a testimoniare una frode elettorale imbastita dalla sinistra. «Non possiamo stimare la diffusione su WhatsApp, in quanto network chiuso. Ma riteniamo che le persone esposte in Brasile alle principali fake news che abbiamo controllato sono nell'ordine delle decine di milioni di elettori», dicono i fact checker brasiliani. «Il dirompente effetto WhatsApp in Brasile non credo sia dovuto ai bassi indici di lettura e scolarità, come qualcuno sostiene – spiega Tai Nalon, giornalista fondatrice di Aos Fatos -. Indagini dimostrano che la pseudonotizia arrivata da un familiare, un amico, la tua chiesa è ritenuta più plausibile in un momento di cattiva reputazione dei media tradizionali, come questo, in tutte le classi sociali». Un altro fattore decisivo sono le offerte degli operatori telefonici. «WhatsApp ormai è gratis, è l'unica cosa che funziona bene dove il segnale è di bassa qualità e molti non possono nemmeno cliccare su un link per controllare la notizia, perché pagherebbero».

Pime azioni

La campagna di Haddad (l'erede di Lula) è riuscita finora ad ottenere dal Tribunale elettorale che vengano rimossi dalla rete i video con i riferimenti al cosiddetto «kit gay», uno dei pezzi da novanta dalla propaganda di Jair Bolsonaro ([che al primo turno ha vinto con quasi il 47% dei voti](#)). Secondo il candidato di estrema destra, negli anni di Haddad come ministro dell'Educazione venne distribuito nelle elementari un libro, «una collezione di assurdità che stimolava precocemente

i bambini ad interessarsi al sesso, alle diverse opzioni sessuali, vale a dire una porta aperta per la pedofilia». La leggenda circola da un paio di anni, è stata smentita numerose volte ma ha continuato a girare in Rete. Da qui la richiesta all'Authority elettorale. Haddad ha anche querelato un noto opinionista di destra che di recente lo ha accusato di essere «a favore dell'incesto». Come si vede, siamo al vale tutto. Ma quella del «kit gay» è stata un'arma potente di diffusione del pensiero di Bolsonaro. Indagini hanno dimostrato che la stragrande maggioranza delle famiglie in Brasile ritiene che l'educazione sessuale è compito esclusivo dei genitori.

Il network chiuso

Quanto all'ipotesi di una azione forte su WhatsApp, non è ancora chiaro se i giudici del Tribunale elettorale accoglieranno le tesi dei loro consiglieri, e ancor meno prevedere che misure concrete si potrebbero adottare a meno di un paio di settimane dal ballottaggio elettorale. Come è noto gli scambi di messaggi su WhatsApp – e app simili – avvengono tra singole utenze telefoniche o tra loro gruppi, e la società ha sempre negato di avere il benché minimo controllo sui contenuti criptati. Molto meno che su Facebook e Twitter, dove testi e foto o gli stessi account possono essere rimossi in caso di reclami di utenti che i sistemi di controllo ritengano fondati. L'unica decisione presa di recente da WhatsApp è indicare quando un messaggio è un forward rispetto ad uno scritto di proprio pugno dalla persona che lo invia. Il problema è ormai universale, ma il Brasile (quarto Paese con più iscritti al mondo, 120 milioni di utenti, 6 volte più che negli Usa) è un caso a sé. Negli ultimi anni, giudici locali sono arrivati a bloccare WhatsApp per alcune ore in tutto il Paese, come rappresaglia contro la casa madre americana che si era rifiutata di fornire dati di utenti utili a indagini criminali.

La “Social Crisis” dei #Ferragnez e il Paradosso del Controllo



Quante aziende oggi si trovano ad affrontare improvvise situazioni di crisi che vengono gestite senza minimamente tenere conto del fatto che ciò che avviene sui social non è controllabile? **Giuseppe Mastromattei**, presidente del Laboratorio per la Sicurezza prende spunto dal recente episodio **Ferragni-Fedez** per un'analisi delle nuove crisi nell'era social.

Nel 1998, usciva nelle sale un Film che ha segnato un'epoca. Mi riferisco a “The Truman Show” diretto da Peter Weir, su soggetto di Andrew Niccol, e interpretato da Jim Carrey, in una delle sue interpretazioni più apprezzate. La pellicola fu

definita allora **una satira fantascientifica**, ispirata parzialmente a un episodio di "Ai confini della realtà" e alla moda allora nascente di raccontare la vita in televisione attraverso i *reality show*, **immaginando una situazione paradossale**, portata all'estremo, dalla quale emergono temi filosofici.

Chi di voi ha visto il film, uscito dalla sala, avrà provato una immensa sensazione di amarezza, tristezza e sconforto, consolato però dal fatto che si trattava di "fantascienza"...

Era il 1998, sono passati 20 anni, sono arrivati gli smartphone i social media, e non è più fantascienza...

Proviamo per un attimo a ritornare indietro nel tempo e analizziamo, con un approccio diverso quello che è successo nei giorni scorsi e che ha visto protagonisti i coniugi Chiara Ferragni e Fedez, ovvero i **Ferragnez** (almeno così mi risulta si facciano chiamare...).

La moglie organizza una festa a sorpresa per il compleanno del marito e sceglie di farla in una location assolutamente innovativa: un supermercato.

Un'idea meravigliosa, quanti di noi hanno desiderato, da bambini, di avere la possibilità di passare una sera in un supermercato ed essere liberi di consumare tutto quello che si vuole.

E fin qui tutto normale, a parte il fatto che trovo leggermente assurda l'idea di poter affittare un supermercato, ma andiamo avanti.

Se fossimo veramente stati nel 1998, forse ci sarebbe stato qualche articolo sui giornali che avrebbe raccontato i fatti, lasciandosi andare a commenti entusiasti per l'idea e la sua originalità, un po' come quella pubblicità di molti anni fa dove un uomo portava al cinema la moglie ed invece di un film, in sala venivano proiettati filmati (in super8) e fotografie, della loro storia d'amore e lui le regalava un diamante. WOW!

Ma purtroppo per i *Ferragnez* siamo nel 2018, e tutto viene condiviso su social, volontariamente; sin dal primo momento, i filmati iniziano nell'ascensore di casa per poi raccontare ogni momento della festa, che ben conosciamo e che trovo

inutile commentare.

Non riesco però a non fare un paragone con un altro film, questa volta un capolavoro di animazione, e soprattutto di soli 10 anni fa: "Wall-e" (nel settembre 2008 Facebook fece il primo salto di visibilità che portò gli utenti, in Italia, a superare il milione e l'iPhone era stato appena lanciato sul mercato).

Bene, nel film c'è una scena dove un uomo e una donna, seppur vicini, si parlano, seduti su una poltrona mobile, attraverso un dispositivo che altro non è uno schermo.

Quello che è successo quella sera ha scatenato una serie infinita di reazioni negative sui social, facendo diventare quella che doveva essere una festa in una vera e propria crisi, una **Social Crisis**, che ha coinvolto non solo la celebre coppia ma anche **Carrefour** che non aveva resistito all'idea di affiancare il proprio marchio a quello dei Ferragnez, sperando in un ritorno in termini di immagine che purtroppo così non è stato, anzi.

Torniamo ai fatti, ad un certo momento si è costituita una unità di crisi, così composta: Chiara Ferragni, Fedez e la madre. Durante una riunione lampo è stato deciso di comunicare, immediatamente che tutto sarebbe stato messo a posto e che sarebbe stata fatta beneficenza, scusandosi di essere stati fraintesi; bella intenzione, forse comunicata troppo in fretta e con un linguaggio non verbale poco adeguato, ma soprattutto senza sapere che sui social andavano in onda i filmati della *"riunione del comitato di crisi"* e del *briefing* fatto prima per decidere i contenuti da diffondere.

Cosa è successo quindi? Ecco il paradosso del controllo, pensando di avere il pieno controllo dell'evento e di quello che si stava svolgendo, sia all'interno del supermercato sia in rete, il controllo è stato perso completamente.

Al punto che una delle prime reazioni è stata quella, dopo l'ennesimo filmato di scuse, questa volta con gli occhi pieni di lacrime, di chiudere uno degli account sui social, proprio come successe al povero Truman, che andando a sbattere sulla scenografia, che delimitava il suo mondo con il mondo reale,

decise di salire le scale ed uscire da una porta di servizio, da cui però non è più rientrato.

Quello che ha fatto una famosa “youtuber” italiana che ha deciso di chiudere il proprio canale (almeno queste le ultime notizie) ed una sua fan ha commentato questa decisione con il seguente *tweet*: “sono contenta che voglia provare a vivere al di fuori di esso...” (sic!)

Aldilà delle considerazioni su questi recenti ed interessanti casi che si sono verificati nei giorni passati, che, proprio come “The Truman Show” fanno emergere infiniti temi filosofici, forse è arrivato il momento di prendere piena consapevolezza dei limiti e dei conseguenti rischi che le nuove tecnologie, ma soprattutto il loro utilizzo eccessivo, possano comportare, non solo per i singoli individui ma anche per le organizzazioni.

Quante aziende oggi si trovano ad affrontare improvvise situazioni di crisi che vengono gestite senza minimamente tenere conto del fatto che ciò che avviene sui social non è controllabile. Proprio per questo motivo è fondamentale tenere presente che in fondo esiste una parete, aldilà della quale c'è il mondo reale, e che basta uscire dalla porta di servizio per rendersi conto che alla fine basta un po' di coerenza e un pizzico di buon senso per non perdere il controllo.

A dimenticavo, “casomai non vi rivedessi, buon pomeriggio, buona sera e buona notte!” (Truman Burbank)

Il conversational commerce nell'era dell'on demand



Il conversational commerce sta cambiando il modo in cui usufruiamo della rete e acquistiamo. Domani essere trovati sul web non sarà più sufficiente, le aziende devono adeguarsi

Viviamo in un mondo in cui tutto è a portata di click. Siamo stati abituati a desiderare e ottenere immediatamente le nostre canzoni e film preferiti attraverso Spotify e [Netflix](#). Siamo ora viziati dalla possibilità di ricevere a casa, nel giro di pochi attimi o di pochissime ore, anche prodotti, cibo, servizi grazie ad Amazon Prime, Deliveroo e tanti altri servizi. È l'economia **on demand**, che dopo aver modificato il rapporto tra persone e aziende **sta ora ristrutturando il ruolo dei canali** di contatto e delle informazioni a disposizione delle persone.

È dagli anni Novanta che troviamo una risposta ai nostri bisogni informativi attraverso la ricerca **sui motori di ricerca** e sui siti web, mutuando peraltro l'approccio che vede l'uomo, da oltre due millenni, informarsi entrando nelle biblioteche e ricercando tra gli scaffali il contenuto più adatto per rispondere ai suoi quesiti.

Con i **social network** abbiamo imparato a raccogliere informazioni più velocemente e ottenere supporto immediato scavalcando i call center e aggirando i menu telefonici (nel dubbio premete sempre il numero 9 e attendete un operatore): al di là dei messaggi più o meno simpatici proposti dai brand

nelle loro pagine, **agli utenti interessa risolvere un problema** e farlo immediatamente, sconvolgendo le linee editoriali e mettendo in crisi i social media manager.

Dopo l'era dello sviluppo del web nei primi anni 2000 e quella della diffusione dei social network e instant messages dieci anni dopo, stiamo ora entrando in un nuovo ciclo legato al **conversational commerce**: una modalità che permette alle persone di **raccogliere informazioni, ottenere supporto e fare shopping** ponendo delle semplici domande sia in modalità testuale, sia attraverso la nostra voce.

Le aziende si stanno adeguando inserendo finestre di chat nei loro siti con operatori disponibili a risponderci immediatamente; in molti hanno previsto chat bot nei loro social network per **offrire supporto 24/7**.

Secondo i dati dello [State of Chatbots Report 2018](#), tra i cui autori figurano aziende come Salesforce, le persone usano i chat bot per avere un supporto continuativo (64% delle risposte), immediato (55%), e ottenere risposte a semplici quesiti (55%). Il contributo dell'uomo **non sarà comunque messo in discussione**: il 43% delle persone preferisce comunque un contatto umano e il 34% dichiara di tollerare un primo contatto con un chat bot, per poi essere connesso a una persona che possa gestire domande complesse e con cui creare un contatto più empatico.

Con l'avvento dei **personal assistant**, accessibili oggi da smartphone e smart speaker, domani da automobili e oggetti connessi, il contatto tra persone, contenuti e prodotti sarà sempre più immediato e l'interfaccia di riferimento sarà la voce.

Da una recente ricerca di Wavemaker è emerso che a **oggi in Italia l'80% degli intervistati** conosce gli assistenti vocali come [Google Assistant](#) o **Siri**, il 20% li utilizza e solo il 5% li considera per la creazione di shopping list in ottica di conversational shopping. Secondo eMarketer, negli Stati Uniti l'utilizzo è del 95%, circa il 40% ricerca vocalmente informazioni sui negozi e il 25% sui prodotti.

Gli smart speaker in particolare stanno **rivoluzionando il**

rapporto tra persone e contenuti: per la prima volta riusciamo a interagire senza prendere in mano un *device* elettronico. eMarketer ha rivisto al rialzo le sue stime e descrive un mercato americano che oggi ha più di 60 milioni tra Amazon Echo (Alexa), Google Home e altri smart speaker. Come spesso accade, guardiamo gli Stati Uniti per ipotizzare cosa potrà succedere anche da noi tra uno o due anni. In una famiglia su quattro oggi si trova uno smart speaker in salotto, cucina o camera da letto, con cui le persone interagiscono costantemente per ascoltare musica (74%), fare domande di varia natura (72%), informarsi su prodotti (39%), attivare oggetti connessi, come per esempio accendere le lampadine in casa (33%) e acquistare prodotti (28%).

È evidente che siamo ancora in una **fase piuttosto embrionale** dell'utilizzo delle voice interface almeno in termini di "conversational shopping". Se ci focalizziamo sul 28% che acquistano prodotti, i "conversational shopper", la maggior parte fa ricerche attraverso gli smart speaker (51%), aggiunge prodotti alla shopping list (36%), chiede informazioni sullo stato di consegna (30%), fa un acquisto vero e proprio (22%), fa una review di prodotto o dà un punteggio (20%).

Le implicazioni per i brand

Per quanto ai primi passi della sua esistenza, il conversational shopping promette di rivoluzionare il modo in cui siamo abituati a concepire la ricerca: oggi cerchiamo una cosa su un motore di ricerca e se è rilevante il fatto di trovarla o meno dipende dall'efficacia del Seo e della Sea.

Domani la ricerca verrà **"filtrata" da un assistente personale**, un'intelligenza artificiale che impara a conoscere le nostre abitudini e i nostri gusti: essere pertinenti da un lato (Seo) e pagare per essere visibili (Sea) semplicemente **non sarà più sufficiente per essere trovati**.

I brand dovranno necessariamente entrare nel repertorio dei consumatori.

Se i **brand** vorranno competere dovranno lavorare nuovamente sul loro posizionamento *top of mind*: i loro prodotti compariranno in cima alle proposte degli smart speaker solo quando saranno chiamati esplicitamente con ricerche del tipo “Alexa, acquista il caffè in cialde della marca ABC” piuttosto che “Alexa, acquista il caffè in cialde”.

Le implicazioni per le aziende sono evidenti: dovranno tornare a lavorare sul *branding*, (ri)creando una forte associazione a un determinato prodotto per vincere le logiche degli algoritmi e apparire in testa alle preferenze.

Il purchase journey

In termini di impatto sul purchase journey il conversational commerce avrà implicazioni tecniche soprattutto nella fase attiva, dove è più rilevante: quando **le persone hanno bisogno** di supporto **occorre rispondere immediatamente**, predisponendo chat bot ma anche customer care (personali o artificiali) che rispondo [attraverso instant messaging](#).

Il purchase journey è un servizio che fotografa il mercato valutando la pubblicità online, quella offline, l'impegno dei marchi o dei brand e il ruolo dei social media nelle esperienze di acquisto dei singoli.

Le **ricerche vocali**, che ora stiamo imparando a fare **attraverso gli smartphone**, devono essere intercettate da chiavi di ricerca che rispecchiano **il linguaggio naturale** e con porzioni di siti web capaci di essere riconosciute e lette dagli assistenti virtuali in modo naturale.

Ma per quanto concerne le implicazioni più strategiche, l'avvento del conversational commerce promette implicazioni ancora maggiori sulla fase del purchase journey che in Wavemaker chiamiamo “priming stage”, quella che precede la ricerca delle informazioni finalizzate all'acquisto.

Noi chiamiamo bias la forza con cui il **priming stage condiziona le scelte di acquisto** ed essa varia a seconda delle categorie oscillando fra un 30% a 60% medio: per convertire l'interesse in vendite occorre sviluppare e mantenere una

forte conoscenza del brand e un **posizionamento chiaro nella mente delle persone.**