

«Accetta e clicca qui» (e la foto di una donna): così il principe saudita ha hackerato Bezos



Mbs avrebbe inviato una fotografia simile alla nuova compagna del proprietario del Washington Post per intimidirlo ma anche per installare uno spyware nel suo telefono

Il messaggio non avrebbe potuto essere più esplicito. L'8 novembre 2018, appena un mese dopo l'assassinio di Jamal Khashoggi, Jeff Bezos, l'uomo più ricco del mondo, riceve un messaggio indesiderato dall'account WhatsApp di Mohammed bin Salman.

Secondo le Nazioni Unite che hanno aperto un'indagine sulla vicenda, il messaggio del principe ereditario dell'Arabia Saudita contiene un'unica fotografia che ritrae una donna bruna. La somiglianza con Lauren Sanchez con cui il miliardario all'epoca ha una relazione clandestina, è evidente. Il messaggio – secondo quanto racconta il *Guardian* – contiene anche un testo che recita: «Litigare con una donna è come leggere il Contrat

«Accetta». Sarebbe iniziato così lo scambio che ha portato all'hackeraggio del telefono di Bezos. Per Agnes Callamard, il relatore speciale delle Nazioni Unite che sta indagando sull'omicidio di Khashoggi, il messaggio è la prova del tentativo da parte della corona saudita di intimidire Bezos . L'obiettivo – questa la teoria – era farlo sentire vulnerabile mentre il suo giornale, il *Washington Post*, continuava a pubblicare storie sull'omicidio di uno dei suoi stessi giornalisti, Jamal Khashoggi, per la cui morte Mbs era già allora il principale indiziato come mandante.

Indietro veloce di qualche mese. Secondo la ricostruzione delle Nazioni Unite, la storia è iniziata il 21 marzo 2018, quando Bezos viene invitato a una piccola cena in onore del principe ereditario la cui lista degli ospiti includeva l'ex giocatore di basket Kobe Bryant e l'amministratore delegato della Disney, Bob Iger. Due settimane dopo, il 4 aprile, i due uomini si scambiano i numeri di telefono a una cena . Il 1 ° maggio, Bezos riceve «un messaggio dall'account del principe ereditario ... tramite WhatsApp», spiega l'Onu. «Il messaggio è un file video crittografato. In seguito viene stabilito, con ragionevole certezza, che il download del video infetta il telefono di Mr Bezos con un codice dannoso». Nei giorni e nelle settimane che seguono, Bezos – che all'epoca era sposato – manda messaggi di testo privati alla sua ragazza, descrivendo i suoi sentimenti. Tali testi saranno successivamente pubblicati dal *National Enquirer*, anche se non è ancora chiaro come il magazine sia entrato in possesso di

questi scambi. Da sottolineare però – fa notare sempre il *Guardian* – come il principe ereditario all'epoca abbia incontrato due volte il proprietario del *National Enquirer*, David Pecker, noto a Hollywood e Washington come un uomo vicino a Donald Trump, la cui presidenza ha rapporti particolarmente stretti con Riad.

Secondo Callamard e Kaye, relatore speciale delle Nazioni Unite sulla libertà di espressione, il «targeting» di Bezos è solo l'inizio di una campagna più ampia per intimidire le persone vicine a Khashoggi e in frequente contatto con il giornalista. La cronologia pubblicata dagli investigatori fa riferimento ad altri quattro importanti oppositori sauditi presi di mira con malware nelle settimane seguenti: Yahya Assiri e Omar Abdulaziz, il comico londinese Ghanem al-Dosari, e un funzionario di Amnesty International che lavorava in Arabia Saudita. Gli investigatori hanno sottolineato anche come ci sia stata «una massiccia campagna online» contro Bezos e Amazon in Arabia Saudita.

Nel rapporto di FTI Consulting – la società che ha analizzato il telefono dell'a.d. di Amazon per conto delle Nazioni Unite – si legge come Bezos abbia avuto un briefing dettagliato sulla campagna saudita contro di lui il 14 febbraio. Poi due giorni dopo, sempre secondo il rapporto FTI, il principe ereditario invia un altro messaggio a Bezos, sostenendo che «ciò che ascolti o dici non è vero ed è giunto il momento che tu dica la verità». L'1° aprile però la campagna contro Bezos cessa. Un fatto da mettere in relazione, probabilmente, con il fatto che Mike Pompeo, il segretario di stato americano, ha negli stessi giorni sollecitato privatamente il principe ereditario a tagliare i suoi legami con il suo stretto consigliere, Saud al-Qahtani, noto come l'uomo della cyber war di Riad.

Sarebbe proprio Al-Qahtani il fautore dell'utilizzo contro Bezos e gli oppositori dello spyware Pegasus, tra i prodotti di punta dalla Nso, società di sicurezza informatica

di Herzliya, in Israele. Pegasus si serve del meccanismo, tutto psicologico, del clickbaiting: invia un messaggio agli utenti di WhatsApp; se una volta aperto il messaggio si clicca sul contenuto, spesso un link, l'azione consente al programma di spia di installarsi sul cellulare, senza che l'utente se ne accorga. Secondo un rapporto di Citizen Lab del 2018 Pegasus negli ultimi anni è diventato molto popolare presso i governi di alcuni paesi del Golfo – Arabia Saudita, Bahrein, Emirati arabi uniti – e per fini anche diversi da quelli del contrasto al crimine. Tra questi: spiare i propri cittadini, per limitarne più efficacemente le libertà e neutralizzare le opposizioni. Pegasus è lo stesso software che all'inizio del 2019 aveva infettato e spiato i cellulari circa 1.400 persone – tra questi molti attivisti politici e giornalisti – attraverso una falla di WhatsApp, popolare sistema di messaggistica di proprietà di Facebook.



[Jeff Bezos✓@JeffBezos](#)

[#Jamal](#)



[27.50020:03 – 22 gen 2020](#) [Informazioni e privacy per gli](#)

[annunci di Twitter8.899 utenti ne stanno parlando](#)

Per il momento tutte le parti in causa negano un coinvolgimento. In un tweet, il governo saudita ha definito «assurde» le accuse. Stessa cosa ha affermato la American Media Inc, proprietaria del National Enquire, che non ha voluto fare ulteriori commenti. L'Arabia Saudita ha insistito sul fatto che il principe ereditario non avesse nulla a che fare con l'omicidio di Khashoggi. Ha anche negato l'uso della tecnologia di sorveglianza contro i critici del regno. Ma, mercoledì sera, mentre la storia continuava a crescere, [Bezos ha postato su Twitter una foto che lo ritrae al funerale di Khashoggi. Come dire, insomma, che la cyber guerra è tutt'altro che finita.](#)

**Quando il cervello fa tilt:
che cosa sono i bias
cognitivi e come vengono
usati contro di noi**



Pregiudizi, fatti sostituiti con impressioni, logiche perfettibili. Quello dei **bias cognitivi** è un campo vasto, queste “défaillance di pensiero” vengono **usate dalla nostra mente per non fare fatica e, basate su percezioni imprecise e pregiudizi** (anche ideologici) ci pongono fuori dal giudizio critico. **Il marketing ne fa incetta** per spingerci a comprare, i movimenti che raccolgono accoliti vi fanno leva per rastrellare seguaci e, come dimostriamo qui, i bias cognitivi rischiano di fare incarcerare innocenti.

Gli **psicologi israeliani Daniel Kahneman e Amos Tversky**, già a partire dagli anni '70 del secolo scorso, hanno studiato a fondo le modalità con cui l'essere umano prende decisioni. Il loro lavoro, che ha contribuito a divulgare il concetto di bias cognitivo (**se ne contano oltre 100**) è stato tuttavia basato su ricerche meno ampie svolte da diversi altri specialisti negli anni precedenti.

Perché sono insidiosi

Iterando gli errori imposti dai bias cognitivi, si rischia di cadere in una forma di pensiero disfunzionale che conduce alla sofferenza emotiva.

Tra i diversi tipi di **bias** cognitivi c'è quello di **conferma**, individuato e teorizzato nei primi anni Cinquanta dallo psicologo americano Burrhus Frederic Skinner secondo il quale **ognuno di noi tenderebbe ad allinearsi a quelle persone o a quelle linee di pensiero che confermano le nostre opinioni**, escludendo così ogni forma di contraddittorio. Così, per esempio, leggiamo soltanto libri o quotidiani che cementano le nostre convinzioni.

Il **bias di gruppo**, molto simile a quello di conferma, induce a **sovrastimare le capacità del gruppo** di cui si fa parte, adducendo a fattori esterni i successi ottenuti da gruppi antagonisti i quali non si riconosce un merito proprio.

A seguire, tra i bias più frequenti, c'è quello di **ancoraggio**. Quando dobbiamo fare una scelta ci **basiamo su fattori ed elementi che riteniamo essere ottimi per fare paragoni** e che, in realtà, non soltanto non sono tali ma ci impediscono di vedere un aspetto nel suo insieme.

Questo tipo di bias è molto usato nelle vendite: **quando si compra uno smartphone da 1.200 euro appare ragionevole comprare un accessorio da 150 euro**, per esempio degli auricolari bluetooth. In realtà, con il prezzo dell'accessorio si può acquistare un cellulare di bassa gamma. **Se acquistiamo una vettura da 100mila euro, spenderne 10mila in optional appare ragionevole** e non è per forza detto che tali optional siano necessari.

Il pregiudizio sullo status quo (**bias sullo status quo**) è molto diffuso. La **situazione attuale è ottimale**, ogni modifica è considerata una regressione. Così **non cambiamo operatore telefonico**, magari a fronte di un'offerta per noi vantaggiosa, perché temiamo costi nascosti o tempi di attivazione lunghi.

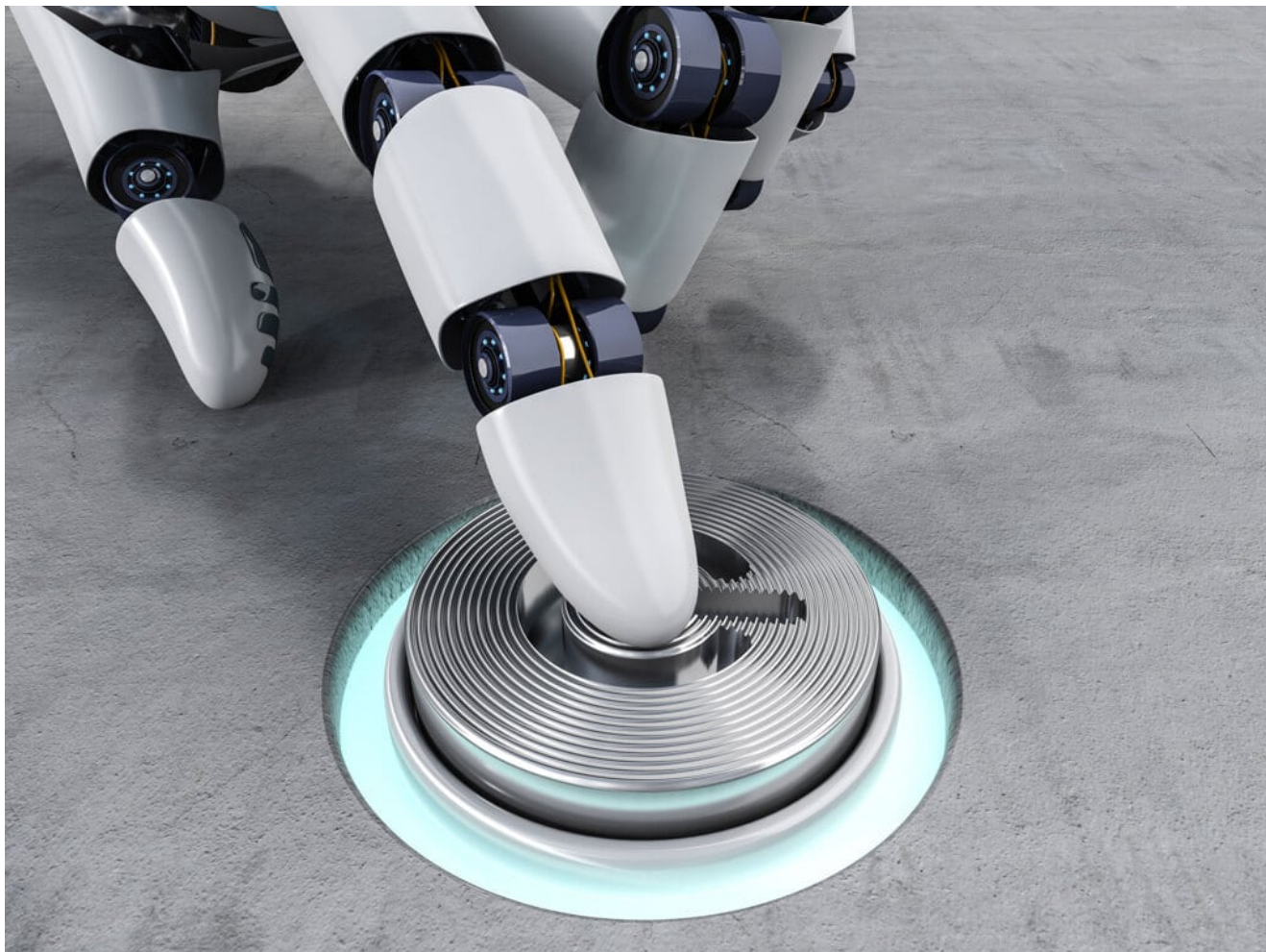
E, semmai ci rendessimo conto di avere preso una decisione perfettibile, il **bias di rinforzo delle scelte** ci viene in soccorso, facendoci ricordare le scelte migliori di quello che

sono state in realtà, mettendone in evidenza (e persino enfatizzando) i punti positivi e **evitando di considerare quelli negativi**.

Il pregiudizio e le sue conseguenze

I bias portano a fare scelte imprecise e non sono sempre scevri di conseguenze pesanti. Lo spiega l'avvocato Carlo Blengino in un articolo pubblicato a [fine 2018 su Il Post](#), in cui parla della **trascrizione di un'intercettazione ambientale**, quindi la trascrizione di un file audio in cui si sentivano le parole "vi ho portati giù ora vi ammazzo". Una frase che non avrebbe lasciato scampo all'imputato se non fosse che, in seguito a un ascolto più attento della stessa traccia audio, il risultato è apparso del tutto diverso: l'imputato ha detto "vi ho portati giù, ora calmatevi". La pessima qualità dell'audio e il **pregiudizio** con cui questo è stato ascoltato, stavano portando le autorità a prendere un granchio.

Il futuro della politica è in mano all'intelligenza artificiale



La stagione della campagna presidenziale statunitense è ufficialmente, ma davvero ufficialmente, arrivata, il che significa che è il momento di affrontare gli strani e insidiosi modi in cui la tecnologia sta distorcendo la politica. Una delle principali minacce che si profilano all'orizzonte è l'arrivo di personalità artificiali, destinate a dominare il dibattito politico. Il rischio nasce da due tendenze che si presentano contemporaneamente: la generazione di testi alimentata dall'intelligenza artificiale e le chatbot (i software che simulano una conversazione con un essere umano) sui social network. Queste "persone" generate da computer sommergeranno le discussioni realmente umane su internet.

I software di generazione di testi sono già abbastanza avanzati da trarre in inganno la maggioranza delle persone, la maggior parte delle volte. Stanno già scrivendo notizie, soprattutto [di sport](#) e [di finanza](#). Parlano con i clienti nei

siti che vendono prodotti. Scrivono [convincenti editoriali](#) su alcuni argomenti d'attualità (anche se esistono [dei limiti](#) al riguardo). E sono usati per rafforzare il “giornalismo pink-slime”, quei siti concepiti per sembrare fornitori di notizie locali ma che in realtà [pubblicano propaganda](#).

Anche i contenuti generati da algoritmi che si presentano come se fossero scritti da esseri umani hanno raggiunto livelli record. Nel 2017, per un certo periodo, la Commissione federale per le comunicazioni degli Stati Uniti ha aperto ai commenti online il suo piano di porre fine alla [neutralità della rete](#), ricevendo l'impressionante cifra di 22 milioni di commenti. Molti di questi, forse la metà, erano falsi, e si servivano d'identità false. Questi commenti erano anche poco elaborati: 1,3 milioni erano [generati a partire dallo stesso modello](#), semplicemente con alcune parole modificate perché apparissero diversi gli uni dagli altri. Non reggevano neanche di fronte a un'analisi sbrigativa.

Disinformazione e democrazia

Simili azioni saranno sempre più sofisticate. Nel corso di un recente esperimento Max Weiss, un ricercatore di Harvard, ha usato un programma di generazione testi per creare mille commenti in risposta a un appello del governo federale relativo al programma sanitario Medicaid. Ciascuno di tali commenti era diverso dall'altro, e sembrava frutto di persone reali che difendevano una posizione politica specifica. Hanno ingannato gli amministratori del sito Medicaid.gov, che li hanno ritenuti reali preoccupazioni di esseri umani in carne e ossa. Trattandosi di ricerca accademica, Weiss ha successivamente identificato i commenti e ha chiesto che fossero rimossi, affinché non vi fosse alcuna interferenza irregolare con l'effettivo dibattito sull'argomento. Il prossimo gruppo che tenterà una cosa del genere non sarà altrettanto onesto.

Sono anni che le chatbot distorcono le discussioni sui social

network. Circa [un quinto dei tweet](#) relativi alle elezioni presidenziali del 2016 è stato pubblicato da bot, secondo una stima. Lo stesso vale per circa [un terzo](#) di quelli relativi al voto sulla Brexit dello stesso anno. [Un rapporto dell'Oxford internet institute](#) dello scorso anno ha trovato prove dell'utilizzo di bot per diffondere propaganda in cinquanta paesi. Questi tendevano a essere programmi semplici che ripetevano automaticamente slogan, come i [250mila tweet filosauditi](#) "abbiamo tutti fiducia in Mohammed bin Salman" apparsi dopo l'[omicidio di Jamal Khashoggi](#) nel 2018.

Il nostro futuro sarà fatto di chiassose discussioni politiche, perlopiù tra bot e altri bot

Individuare molti bot con pochi follower è più difficile che rilevare alcuni bot con molti follower. E misurare l'efficacia di questi bot non è semplice. Le [migliori analisi](#) indicano che questi non hanno influenzato le elezioni presidenziali statunitensi del 2016. Più probabilmente distorcono la percezione che le persone hanno dell'opinione pubblica e la loro fiducia nella discussione politica ragionata. Siamo tutti immersi in un nuovo esperimento sociale.

Nel corso degli anni, i bot algoritmici si sono evoluti [fino ad avere una propria personalità](#). Possiedono nomi falsi, false biografie e false foto, talvolta generate dall'intelligenza artificiale. Invece di diffondere propaganda senza sosta, postano solo occasionalmente. I ricercatori possono rilevare che si tratta di bot e non persone in base alla frequenza e all'andamento dei loro post. Ma la tecnologia dei bot continua a migliorare, rendendo difficili i tentativi di rilevamento. I gruppi futuri non saranno così facili da identificare. Riusciranno a integrarsi meglio nei gruppi sociali di persone in carne e ossa. La loro propaganda sarà più sottile, e si mescolerà all'interno delle discussioni che interessano tali gruppi.

Mettete insieme queste due tendenze e avrete la ricetta per fare in modo che le discussioni non umane prendano il sopravvento sulle discussioni politiche tra esseri umani.

Chi controlla i bot

Presto le personalità alimentate da intelligenza artificiale saranno in grado di scrivere lettere personalizzate a giornali e parlamentari, esprimere il proprio commento nel quadro di processi legislativi pubblici, creando personalità che perdurano e che appaiono reali anche a quanti cercano di smascherarle. Saranno anche in grado di presentarsi come individui sui social network e d'inviare testi personalizzati. Avranno milioni di repliche e discuteranno di tali questioni giorno e notte, inviando miliardi di messaggi, lunghi e brevi. La somma di tutte queste cose gli permetterà di sommergere ogni reale discussione su internet. Non solo sui social network, ma dovunque ci sarà un dibattito.

Magari questi bot dotati di personalità saranno controllati da attori stranieri. Magari da gruppi politici nazionali. Magari dai candidati stessi. Più probabilmente, chiunque potrà farlo. La più importante lezione a proposito della disinformazione nel 2016 non è che ci sia stata disinformazione, bensì quanto sia stato facile e poco costoso disinformare le persone. I futuri miglioramenti della tecnologia renderanno la cosa ancora più economica.

Il nostro futuro sarà fatto di chiassose discussioni politiche, perlopiù tra bot e altri bot. Non è quello che si ha in mente quando si loda il mercato delle idee, o qualsiasi altro processo politico democratico. La democrazia ha bisogno di due cose per funzionare efficacemente: informazione e rappresentanza. Le personalità artificiali possono privare le persone di entrambe le cose.

Sistemi di difesa

È difficile immaginare delle soluzioni. Possiamo regolamentare l'uso dei bot – una [proposta di legge in California](#) imporrebbe ai bot d'identificarsi – ma la cosa sarebbe efficace solo con le campagne d'influenza legittime, come la pubblicità. Sarà molto più difficile rilevare le operazioni d'influenza surrettizia. La difesa più ovvia è sviluppare e standardizzare metodi migliori di autenticazione. Se i social network verificano che dietro ogni account c'è effettivamente una persona reale, allora saranno in grado di eliminare più facilmente le personalità false. Ma account falsi sono già regolarmente creati per persone reali senza che queste lo sappiano o vi acconsentano, e le discussioni anonime sono essenziali per un sano dibattito politico, soprattutto quando chi parla proviene da comunità marginalizzate o penalizzate. Non abbiamo un sistema di autenticazione in grado al contempo di proteggere la privacy e che sia efficace per miliardi di utenti.

Possiamo sperare che la nostra capacità d'identificare le personalità artificiali tenga il passo con la nostra capacità di mascherarle. Se la lotta sempre più feroce tra *deepfake* e rilevatori di *deepfake* può fungere da guida, anche questo non sarà un compito facile. Le tecnologie di offuscamento sembrano sempre un passo avanti alle tecnologie di rilevamento. E le “persone” artificiali saranno progettate per agire esattamente come le persone reali.

In ultima istanza le soluzioni dovranno essere di natura non tecnologica. Dobbiamo riconoscere i limiti del dibattito politico in rete, e dare nuovamente priorità alle interazioni di persona, che sono più difficili da automatizzare e ci permettono di sapere che le persone con cui parliamo sono esseri umani in carne e ossa. Sarebbe una svolta culturale che permetterebbe di prendere le distanze dai testi pubblicati su internet, e di tenersi lontani dai social network e dai thread

di commento.

I tentativi di disinformazione sono ormai diffusi in tutto il mondo, e sono praticati in più di settanta paesi. È il metodo ordinario con cui si effettua la propaganda nei paesi con tendenze autoritarie, e sta diventando il modo per portare avanti una campagna politica, che si tratti di un candidato o di una questione specifica.

Le personalità artificiali sono il futuro della propaganda. E anche se potrebbero non essere in grado di spostare il dibattito politico in una direzione o in un'altra, possono facilmente sommergerlo del tutto. Non sappiamo quali siano gli effetti di una simile interferenza sulla democrazia, se non che è nociva, e inevitabile.

**Louis Vuitton veste i
personaggi di League of
Legends**



La moda sbarca nel mondo del gaming. Louis Vuitton ha annunciato l'avvio di una partnership con Riot Games, la casa di videogiochi di League of Legends. Una partnership a più livelli con uno sguardo rivolto al futuro, come racconta Sara Bassi.

Lo scorso 23 settembre Louis Vuitton ha annunciato l'avvio di una partnership con Riot Games, la casa di videogiochi famosa per aver dato i natali al gioco online *League of Legends*. La collaborazione prevede diversi livelli, cominciando dalla realizzazione del baule contenente il trofeo del campionato mondiale del videogioco, le cui finali si sono tenute il 10 novembre a Parigi.

La Maison francese si è già resa protagonista di creazioni simili, come la custodia della FIFA World Cup e la Rugby World Cup, ma è la prima volta che avviene per un torneo di eSports. Il design del baule presenta elementi classici della casa di moda ed elementi innovativi ispirati al videogioco. Il risultato è un mix tra la texture con il monogramma, le

rifiniture tipiche delle valigie firmate LV e un tocco high-tech.

Naz Aletaha, l'Head of Global eSports Partnership di Riot Games ha dichiarato di essere entusiasta per la collaborazione con LV perché "i suoi motivi influenzano l'aspetto, l'atmosfera e il prestigio del nostro evento di League of Legends". La pensa così anche il CEO della casa di moda parigina, sottolineando la presenza del marchio accanto ai trofei più ambiti del mondo e come ormai anche la Summoner's Cup (questo il nome del trofeo) lo sia diventata.

La partnership ha portato anche alla creazione di skin (vesti grafiche) per i personaggi di Qiyana e Senna disponibili per un tempo limitato, di cui una è stata presentata proprio in occasione del League of Legends World Championship, e l'altra sarà rivelata il prossimo febbraio. A firmare questi outfit virtuali è stato Nicolas Ghesquière, direttore artistico della linea donna di Louis Vuitton. A differenza di quanto si può immaginare, queste skin non avranno un costo eccessivo e per comprarle basterà guadagnare punti giocando alcune missioni speciali.

Più recentemente, invece, la casa di moda ha lanciato la linea d'abbigliamento ispirata al videogioco. La collezione limited edition è venduta nel negozio online presente sul sito e in alcuni store in giro per il mondo. Anche stavolta il design riesce a coniugare bene l'immagine del marchio LV con lo stile fantasy-futuristico del videogioco.

Non si tratta, però, della prima volta che il mondo videoludico entra nel mondo di Louis Vuitton. Già qualche anno fa, lo stesso Ghesquière, aveva collaborato con la casa Square Enix, per avere come modella per una campagna Lightning, una dei protagonisti della storica saga di Final Fantasy. Il direttore artistico, infatti, sembra avere un debole per la fantascienza e la tecnologia, che cerca sempre di portare all'interno delle sue collezioni.

La distanza tra il mondo della moda e il mondo del gaming sembra ridursi sempre di più. Questo tipo di partnership è sicuramente innovativa e introduce un lato della moda che potrebbe prendere piede nel prossimo futuro, quello dei vestiti virtuali. I videogiochi stanno altresì crescendo sempre di più nella loro diffusione e fruizione, anche da parte di adulti e famiglie. Inoltre, i giocatori sono ormai abituati agli acquisti all'interno del gioco e questo è un aspetto da non sottovalutare.

Unilver VS obesità infantile



L'obesità, soprattutto quella infantile, sta diventando un problema sempre più attuale nei paesi più industrializzati del pianeta. L'OMS, in un comunicato reso pubblico l'ottobre scorso, stimava che entro il 2030 i bambini

e i ragazzi obesi saranno 254 milioni, il 60% in più rispetto a quelli del 2016, se non si interviene immediatamente. A questo proposito alcuni governi hanno iniziato a tassare i cibi contenenti un elevato quantitativo di zuccheri, considerati più dannosi per la salute. Se in Italia solo con la legge di bilancio 2020 è stata introdotta la cosiddetta "sugar tax", cioè la tassa sulle bibite zuccherate, nel Regno Unito una tassa simile è già in vigore dal 2016. Nella fattispecie italiana la legge prevede una tassa pari al 10 centesimi a litro per le bevande zuccherate ad eccezione delle bevande dolcificate con meno di 25 grammi di zucchero per litro. La tassa è stata introdotta sia per disincentivare l'uso di bevande unhealthy ma anche nella speranza che si segua l'esempio del Regno Unito dove molti produttori hanno dimezzato le quantità di zuccheri pur di non far rientrare i loro prodotti tra quelli soggetti a tassazione.

In linea con gli intenti dei governi il colosso inglese Unilever, che opera nel campo dell'alimentazione e delle bevande, ha deciso di modificare considerevolmente la propria strategia di marketing. L'azienda infatti ha deciso di smettere di rivolgere le proprie pubblicità di cibo e bevande ai bambini a partire dal 2020. Unilever non è nuova a questo tipo di iniziative, già nel 2003 aveva pubblicato una lista di principi guida per la

comunicazione dei suoi prodotti di food and beverage, escludendo la promozione di prodotti che contrastassero con una dieta salutare ed equilibrata. Nel 2010 l'azienda aveva firmato, insieme ad altri 18 leader del settore alimentare, la ''children's food and beverage advertising initiative'', dichiarando che, ai bambini sotto gli 11 anni, avrebbe rivolto soltanto pubblicità di prodotti che rispettassero degli standard nutrizionali.

La proposta attuale fa un ulteriore passo avanti con l'azienda che decide di rivolgere la sua comunicazione non più ai bambini, che possono essere facilmente suggestionabili, ma ai genitori che sono quindi in grado di valutare se un prodotto può essere più o meno adatto. Questa modifica non riguarda solo le pubblicità nei canali tradizionali, come gli annunci televisivi o cartacei ma anche tutte le attività online, sui social media e sulle app. L'azienda, tra le altre cose, ha annunciato che non saranno più scelti, per pubblicizzare i prodotti, influencer che abbiano come target principale i bambini sotto i 12 anni.

In particolare, per quel che riguarda il settore dei gelati (Unilever possiede il brand Algida), l'azienda dichiara che le comunicazioni avverranno in maniera responsabile e che i gelati rivolti ai bambini, entro la fine del 2020, non avranno più di 110 calorie e conterranno massimo 12 grammi di zucchero a porzione. Su

questi prodotti, per renderli più facilmente riconoscibili, verrà applicato il logo di "responsibly made for kids" al fine di aiutare i genitori ad identificarli.