

Sommario

CAPITOLO I. INTRODUZIONE	3
1.1 INTELLIGENZA ARTIFICIALE E COMUNICAZIONE DI MARKETING	4
1.2 OOH E DOOH: UN MEZZO TRADIZIONALE DIGITALIZZATO	5
1.3 CARENZE NELLA LETTERATURA SU AI X DOOH	5
1.4 OBIETTIVO DELLA TESI, DOMANDE DI RICERCA E METODOLOGIA	6
1.5 CONTRIBUTI TEORICI E IMPLICAZIONI MANAGERIALI	7
1.6 STRUTTURA DELLA TESI	7
CAPITOLO II. REVISIONE TEORICA E DELLA LETTERATURA	8
2.1 DEFINIRE L'INTELLIGENZA ARTIFICIALE E LA GENERATIVE AI NEL MARKETING	8
2.2 AI COME STRUMENTO DI COMUNICAZIONE DI MARKETING: COSA STUDIA LA LETTERATURA E DOVE SI CONCENTRA	9
2.2.1 <i>AI come infrastruttura data-driven per pianificare, targettizzare e ottimizzare la comunicazione</i>	10
2.2.2 <i>AI come motore creativo e strumento di produzione e trasformazione dei contenuti</i>	11
2.2.3 <i>Una ricerca sbilanciata verso i media digitali</i>	12
2.2.4 <i>Oltre il digitale: dove la letteratura esiste ma è più frammentata</i>	13
2.3 OOH E DOOH NELLA LETTERATURA: CARATTERISTICHE DEL MEZZO, DETERMINANTI DI EFFICACIA E IMPLICAZIONI TEORICHE	14
2.3.1 <i>OOH come mezzo pubblicitario: attenzione, memoria e contesto</i>	14
2.3.2 <i>La frammentazione della letteratura OOH</i>	15
2.3.3 <i>DOOH: natura del mezzo, evidenze di letteratura e opportunità di ricerca</i>	15
2.4 AI E ADVERTISING: ORIGINE CREATIVA E DISCLOSURE COME SEGNALI PERSUASIVI	18
2.4.1 <i>Origine creativa dei contenuti: "AI-generated" vs "human-created"</i>	18
2.4.2 <i>Disclosure dell'uso dell'AI: trasparenza, credibilità e effetti negativi</i>	19
2.4.3 <i>Origine creativa, disclosure e implicazioni per il DOOH</i>	20
2.5 CORNICI TEORICHE PER SPIEGARE GLI EFFETTI NEL DOOH: PERSUASIONE, SEGNALI, CONOSCENZA DELLA PERSUASIONE E REACTANCE	20
2.5.1 <i>Elaboration Likelihood Model: Dual-process e vincoli di elaborazione in condizioni di "fast exposure"</i>	21
2.5.2 <i>Signaling theory: disclosure e origine creativa come segnali strategici</i>	21
2.5.3 <i>Persuasion Knowledge Model: riconoscimento dell'intento persuasivo</i>	22
2.5.4 <i>Reactance theory: minaccia alla libertà, intrusività e backlash</i>	22
2.6 ELEMENTI CHIAVE DEL MODELLO: DEFINIZIONI, RILEVANZA E STATO DELLA RICERCA	22
2.6.1 <i>Attitude toward the Ad (Aad)</i>	23
2.6.2 <i>Memoria di brand: brand recognition</i>	23
2.6.3 <i>Autenticità, fiducia e credibilità</i>	24
2.6.4 <i>Persuasion knowledge, reactance e resistenza persuasiva</i>	24
2.6.5 <i>Pertinenza e utilità</i>	25
2.7 DALLA LETTERATURA AL GAP	25
2.8 DAL GAP ALLE DOMANDE DI RICERCA	26
2.9 SVILUPPO DELLE IPOTESI DI RICERCA	26
CAPITOLO III. METODOLOGIA	30
3.1 INTRODUZIONE AL CAPITOLO: OBIETTIVO DELLA RICERCA E APPROCCIO METODOLOGICO	30
3.2 DISEGNO DI RICERCA	30
3.2.1 <i>Struttura complessiva</i>	30
3.2.2 <i>Tipologia dello studio e giustificazione dell'approccio sperimentale</i>	31
3.2.3 <i>Variabili del modello e ruolo nel framework di ricerca</i>	32
3.3 PARTECIPANTI E CAMPIONAMENTO	33
3.4 STRUMENTI DI RICERCA E RACCOLTA DATI: PIATTAFORMA, QUESTIONARIO E STIMOLI	34
3.5 PROCEDURA EMPIRICA: STUDIO 1	36
3.5.1 <i>Disegno sperimentale, manipolazioni e randomizzazione</i>	36
3.5.2 <i>Stimoli dell'esperimento e manipolazioni</i>	37
3.5.3 <i>Struttura del questionario</i>	38

3.5.4 Misure e scale utilizzate.....	39
3.6 PROCEDURA EMPIRICA: STUDIO 2	41
3.6.1 Disegno sperimentale, manipolazioni e randomizzazione	41
3.6.2 Stimoli dell'esperimento e manipolazioni	42
3.6.3 Struttura del questionario	43
3.6.4 Misure e scale utilizzate.....	44
3.7 ANALISI DEI DATI	45
3.7.1 Struttura e preparazione del dataset	46
3.7.2 Codifica delle variabili e costruzione degli indici	47
3.7.3 Analisi preliminari	47
3.7.4 Strategia di analisi inferenziale e test delle ipotesi	48
3.7.5 Analisi Studio 1: test delle ipotesi H1–H5	48
3.7.6 Analisi Studio 2: ipotesi H6–H9	51
CAPITOLO IV. RISULTATI	53
4.1 RISULTATI ANALISI PRELIMINARI – STUDIO 1	53
4.2 RISULTATI TEST D'IPOTESI – STUDIO 1	53
4.3 RISULTATI ANALISI PRELIMINARI – STUDIO 2	56
4.4 RISULTATI TEST D'IPOTESI – STUDIO 2	56
CAPITOLO V. DISCUSSIONE E CONCLUSIONI	58
5.1 DISCUSSIONE E INTERPRETAZIONE DEI RISULTATI	58
5.2 CONTRIBUTI TEORICI DELLA RICERCA	59
5.3 IMPLICAZIONI MANAGERIALI DELLA RICERCA	60
5.4 LIMITAZIONI DELLA RICERCA	61
5.5 DIREZIONI DI RICERCA FUTURE	62
5.6 CONCLUSIONI FINALI	63
BIBLIOGRAFIA	65

Capitolo I. Introduzione

Negli ultimi anni, la comunicazione di marketing ha assistito ad una tendenza: lo spostamento dell'attenzione verso i nuovi media digitali. La pianificazione della comunicazione di marketing, infatti, si è spostata verso i media digitali, più adatti per misurare più facilmente l'esposizione, targetizzare audience specifiche e ottimizzare le campagne quasi in tempo reale. Questo cambiamento ha contribuito a ridimensionare il ruolo di diversi media tradizionali, considerati più difficili da gestire rispetto ai canali digitali. In questo contesto, l'Out-of-Home (OOH) ha risentito in modo particolare di questo cambiamento. Il mezzo, infatti, nonostante la sua capacità di copertura nello spazio fisico e la possibilità di raggiungere un ampio pubblico, è collegato a logiche difficili da integrare con i nuovi modelli tipici del digitale (Wilson, 2023).

Con l'evoluzione in Digital Out-of-Home (DOOH), tuttavia, l'OOH ha acquisito caratteristiche più vicine ai nuovi media, diventando più adatto ad essere integrato nelle nuove strategie di marketing mix. Il mezzo sta, infatti, vivendo una fase di rinnovamento grazie a questa sua evoluzione digitale, che gli ha permesso di implementare schermi digitali, contenuti dinamici, possibilità di pianificazione più flessibile e integrazione di dati di contesto per adattare i messaggi all'ambiente. Questi elementi suggeriscono che, in questa epoca dominata dai mezzi digitali, l'OOH non è scomparso, ma si sta trasformando.

Dall'altra parte, l'intelligenza artificiale (AI) sta trasformando tutti i processi della comunicazione di marketing, portando velocità, ottimizzazione e personalizzazione dei contenuti (Davenport et al., 2020; Grewal et al., 2025). Accanto alle opportunità, però, emergono anche dei problemi collegati all'AI, perché il suo utilizzo nella creazione e gestione dei contenuti ha portato a delle regolamentazioni in termini di trasparenza, introducendo degli obblighi in questo senso: informare quando un contenuto sia stato generato o alterato da sistemi di AI, con l'obiettivo di ridurre i rischi di inganno e manipolazione. In termini di comunicazione di marketing, ciò significa che la disclosure, ovvero la dichiarazione esplicita di aver utilizzato AI, non è più solo una scelta, ma diventa un vero e proprio obbligo. La disclosure, però, rappresenta un punto sensibile, dal momento che molte evidenze mostrano che dichiarare esplicitamente l'uso di AI può influenzare negativamente la credibilità percepita e le valutazioni dell'annuncio (Baek et al., 2024).

La presente tesi si colloca proprio all'interno di questo delicato scenario. L'obiettivo specifico del lavoro è, infatti, studiare se e come l'utilizzo di strumenti AI in un contesto DOOH possa aumentarne

l'efficacia pubblicitaria, riuscendo a posizionare il mezzo in primo piano all'interno delle strategie di comunicazione e contribuendo così ad una "rinascita" dell'OOH. In particolare, verrà misurato in modo causale l'impatto dell'AI sull'efficacia pubblicitaria del DOOH, chiarendo i meccanismi che spiegano quando questo impatto diventa positivo e quando, invece, può essere neutro o negativo.

1.1 Intelligenza artificiale e comunicazione di marketing

Nella letteratura scientifica, l'AI viene spesso descritta come un insieme di tecnologie che possono modificare sia le strategie delle imprese sia i comportamenti dei consumatori, intervenendo su varie attività diverse, che possono essere di tipo analitico, relazionale o creativo (Davenport et al., 2020). Con l'avvento della Generative AI, la dimensione creativa ha assunto una rilevanza particolare, portando la tecnologia dall'essere un semplice strumento di supporto, all'entrare direttamente nel pieno delle attività di comunicazione, producendo testi, immagini e video in modo rapido e coerente. Tuttavia, l'adozione dell'AI in comunicazione non implica automaticamente un miglioramento dell'efficacia. La letteratura esistente in materia evidenzia che l'AI è percepita dai consumatori in maniera più complessa, generando valore ma anche creando preoccupazioni e diffidenza. Puntoni et al. (2021) propongono una prospettiva esperienziale dell'Intelligenza Artificiale, affermando che i consumatori possono percepire l'AI sia come beneficio che come costo psicologico; per questo, anche quando l'AI produce risultati oggettivamente efficaci, i consumatori possono comunque risultare diffidenti.

Nel campo dell'advertising, inoltre, un tema delicato è quello della trasparenza, perché, come già detto, la disclosure dell'uso di AI può influenzare la credibilità percepita del contenuto e l'atteggiamento verso l'annuncio, attivando meccanismi di resistenza e scetticismo. Questo non significa, però, che la presenza di disclosure sia sempre un elemento negativo, ma suggerisce che la presenza dell'etichetta può costituire un segnale, che può spingere il consumatore ad attribuire minore contributo umano e minore autenticità e a temere manipolazioni.

1.2 OOH e DOOH: un mezzo tradizionale digitalizzato

La letteratura esistente riguardo OOH e DOOH sottolinea due caratteristiche particolari del media. In primo luogo, l'OOH è descritto come un mezzo incidentale: il consumatore non sempre sceglie di esporsi all'annuncio, ma lo incontra mentre svolge altre attività quotidiane, ponendo i contenuti DOOH a competere con altri elementi, in un ambiente ricco di stimoli e con risorse attentive degli utenti limitate. In secondo luogo, nonostante sia molto diffuso nella pratica, l'OOH è tendenzialmente un ambito di ricerca sottosviluppato e frammentato. In questo contesto, Wilson (2023) ha sviluppato una systematic review in cui consolida oltre un secolo di studi e nella quale sottolinea la dispersione della ricerca e la necessità di ottenere maggiori evidenze causali. In questo contesto, l'utilizzo di AI nel DOOH rappresenta un'area di ricerca in cui i contributi sono ancora più limitati, nonostante la crescente importanza manageriale.

Il DOOH rende il mezzo tradizionale più vicino alla logica e alle caratteristiche dei media digitali, introducendo contenuti dinamici, possibilità di avere rotazioni creative, potenziale programmaticità dei contenuti e adattamento contestuale. Questi elementi suggeriscono che il DOOH può essere più di un semplice cartellone urbano digitale e che può diventare un mezzo rilevante nel contesto. L'AI può inserirsi in questo scenario come un acceleratore delle caratteristiche digitali del DOOH, permettendo di produrre varianti creative rapidamente, abilitando l'adattamento del messaggio al contesto e rafforzando la possibilità di ottimizzazione. Ma, come l'AI può rendere il DOOH ancora più digitalizzato, allo stesso modo può anche introdurre i problemi tipici del digitale, come sfiducia da parte degli utenti e percezione di sorveglianza e controllo.

1.3 Carenze nella letteratura su AI x DOOH

Nella letteratura esistente, esistono varie evidenze riguardo all'AI in contesti di marketing in riferimento a canali e media digitali, che rappresentano ambienti con grande disponibilità di dati e tendenzialmente facili da sperimentare. Dall'altra parte, la letteratura su OOH e DOOH in questo senso resta più frammentata e meno centrata sulle applicazioni dei nuovi strumenti tecnologici.

Questo crea un dualismo teorico ed empirico: da un lato, l'AI ha delle caratteristiche che potrebbero rendere il DOOH più efficace e quindi più competitivo rispetto ai nuovi media; dall'altro, non ci sono

sufficienti evidenze che studino come i consumatori reagiscano a un contenuto DOOH creato o ottimizzato attraverso l'utilizzo di AI.

Per questo, la presente tesi si propone fornire evidenze causali sull'effetto dell'utilizzo di AI nel contesto di DOOH, sfruttando due leve strategiche:

- 1) Origine creativa dei contenuti: se la Generative AI può aiutare nella creazione di contenuti, è necessario capire se, nel DOOH, questa origine della creatività produce vantaggi in termini di efficacia oppure se, al contrario, attiva dei meccanismi negativi che ne riducono la qualità percepita;
- 2) Pertinenza context-aware: se l'AI consente di adattare i messaggi al contesto esterno, da una parte può aumentare la rilevanza e l'utilità percepita dei contenuti, dall'altra potrebbe essere interpretata come troppo intrusiva.

Per entrambe queste leve, la presenza o assenza di disclosure rappresenta un elemento chiave, dal momento che modifica la percezione del messaggio, potenzialmente influenzandone la credibilità percepita.

1.4 Obiettivo della tesi, domande di ricerca e metodologia

L'obiettivo generale della tesi è studiare come l'AI possa contribuire al riposizionamento del DOOH al centro delle strategie di comunicazione, studiando la capacità del mezzo di generare buoni risultati in termini di efficacia pubblicitaria e renderlo potenzialmente in grado di competere in un'epoca dominata dai media digitali.

Per raggiungere questo obiettivo, la tesi si propone di rispondere a tre domande di ricerca. La prima tratta l'efficacia dei contenuti DOOH la cui origine creativa è AI-generated. La seconda è relativa alla disclosure e alla possibilità che questa possa diminuire o amplificare l'efficacia della creatività AI-generated. La terza, infine, si basa sulla pertinenza context-aware e sui suoi effetti di efficacia. Queste domande servono a dimostrare quale combinazione di elementi e quali condizioni generano valore ed efficacia e quali, invece, rischiano di avere conseguenze negative.

Per rispondere alle domande di ricerca, la tesi utilizza una metodologia di ricerca quantitativa, condotta attraverso l'implementazione di due studi sperimentali volti a stimare gli effetti causali e a testare le varie interazioni del modello.

1.5 Contributi teorici e implicazioni manageriali

La tesi mira a produrre sia contributi teorici che implicazioni manageriali.

Per quanto riguarda la teoria, la ricerca estende la letteratura esistente su AI e comunicazione di marketing a un contesto relativamente meno esplorato, il DOOH. In particolare, collega l'AI a outcome classici dell'efficacia pubblicitaria, integrando meccanismi considerati fondamentali quando si parla di AI. In questo modo, contribuisce anche alla letteratura su OOH e DOOH, andando a fornire evidenze sperimentali sistematiche sul ruolo delle nuove tecnologie.

Dal punto di vista manageriale, i risultati della tesi puntano ad offrire implicazioni riguardo alle opportunità e ai limiti della creatività AI-generated nel contesto di DOOH, alle conseguenze strategiche della presenza o meno di disclosure dell'uso di AI e alle possibilità della pertinenza context-aware come leva di rilevanza.

1.6 Struttura della tesi

La tesi da qui in poi è strutturata in quattro capitoli. Il Capitolo II sviluppa il quadro teorico e la literature review alla base dello studio di ricerca, includendo le cornici teoriche dell'AI nel contesto di marketing, le evidenze sulla presenza di disclosure e le relative reazioni dei consumatori, la letteratura esistente su OOH e DOOH e, infine, l'identificazione del gap e la formulazione delle domande di ricerca e delle ipotesi del modello. Il Capitolo III descrive la metodologia utilizzata nel corso della ricerca empirica e nella conduzione dei due studi, definendo stimoli, design, campione, procedure, misure e piano di analisi degli esperimenti. Il Capitolo IV presenta i risultati dei due esperimenti, relativamente agli effetti diretti principali, alle interazioni tra le varie componenti del modello e ai test di mediazione. Infine, il Capitolo V discute i risultati ottenuti, evidenziando le implicazioni teoriche, i contributi manageriali, i limiti dello studio e le possibili direzioni di ricerca futura.

Capitolo II. Revisione teorica e della letteratura

La presente tesi si colloca nel mezzo di due diversi filoni di ricerca che non sono mai stati integrati tra di loro in modo univoco: da un lato, la ricerca su AI e generative AI applicate alle attività di marketing e di comunicazione; dall'altro, la ricerca su OOH e DOOH come mezzo pubblicitario. La combinazione di questi due filoni è teoricamente promettente poiché l'utilizzo di AI può trasformare l'OOH in un media più dinamico e data-driven, aumentandone l'efficacia; dall'altra parte, però, è anche potenzialmente problematica dal momento che l'applicazione di AI può generare timori e diffidenza, con ricadute sulla risposta dei consumatori.

2.1 Definire l'Intelligenza Artificiale e la Generative AI nel marketing

Nel panorama accademico attuale, l'Intelligenza Artificiale non è descritta e trattata come una semplice tecnologia, ma come un insieme di metodi e sistemi in grado di svolgere compiti che normalmente richiedono capacità umane, grazie all'utilizzo di algoritmi e modelli computazionali (Davenport et al., 2020; Huang & Rust, 2021). Nel contesto del marketing, l'AI è rilevante poiché è in grado di inserirsi all'interno dei processi di marketing, modificandone la logica di progettazione e l'implementazione (Davenport et al., 2020).

Una distinzione da sottolineare, particolarmente rilevante per la presente tesi, è quella tra AI analitico-predittiva e Generative AI (GenAI). La prima include quei sistemi AI capaci di classificare e prevedere ed è tendenzialmente utilizzato in contesti di marketing analytics e di ottimizzazione dell'advertising (Davenport et al., 2020). La Generative AI, invece, include quegli strumenti AI in grado di creare contenuti originali, cambiando il ruolo dell'AI e ponendola al servizio di alcuni tra gli elementi più importante dell'advertising: la creatività, l'esecuzione e lo stile del messaggio (Grewal et al., 2025). Nel contesto specifico della comunicazione di marketing, la duplice natura dell'Intelligenza Artificiale in AI analitico-predittiva e Generative AI implica che l'AI possa intervenire sia sulla costruzione della delivery dei contenuti, attraverso l'ottimizzazione di targeting, frequenza, timing e placement di esposizione, che sulla creazione dei contenuti stessi, attraverso la produzione e l'adattamento della creatività.

Questa distinzione è spesso collegata alla differenza tra automation e augmentation. In questa prospettiva, l'AI può sostituire alcune componenti del lavoro umano (automation), ma può anche

potenziare la capacità umana (augmentation) migliorandone rapidità, precisione o creatività (Huang & Rust, 2021).

La letteratura evidenzia poi che l'adozione e l'utilizzo di AI nella comunicazione di marketing non produce solo effetti tecnici e operativi, ma genera anche effetti psicologici nella mente dei consumatori e nelle loro percezioni. Puntoni et al. (2021) propongono una prospettiva basata sulle esperienze, secondo la quale i consumatori non interagiscono con l'AI come se fosse uno strumento neutro, ma come un insieme di elementi che possono generare sia benefici che costi psicologici. In particolare, quando l'AI diventa in qualche modo esplicitamente evidente può attivare nella mente di consumatori pensieri riguardo al modo in cui l'azienda opera e alle sue intenzioni persuasive, con potenziali conseguenze negative su credibilità e autenticità percepita e aumento di resistenza verso gli annunci (Puntoni et al., 2021).

2.2 AI come strumento di comunicazione di marketing: cosa studia la letteratura e dove si concentra

Negli ultimi anni, la letteratura sul marketing ha descritto l'Intelligenza Artificiale come un insieme di tecnologie, in grado di supportare o automatizzare decisioni e attività di marketing, inclusa la comunicazione. L'intelligenza artificiale, però, non è solo un insieme di tecnologie di supporto, ma è una vera e propria forza capace di trasformare processi, contenuti e relazioni tra brand e consumatori. In questa prospettiva, l'AI diventa uno strumento comunicativo che interviene lungo l'intero processo di comunicazione, influenzando ciò che il consumatore vede, quando lo vede e in quale contesto (Huang & Rust, 2021; Grewal et al., 2025). L'AI, infatti, può essere utilizzata lungo l'intero customer journey, influenzando la raccolta e l'elaborazione dei dati, le decisioni di targeting e delivery dei messaggi, la produzione e l'adattamento creativo dei contenuti e le modalità con cui il consumatore interagisce e si espone alla comunicazione.

Alcuni modelli ricorrenti nella letteratura propongono delle classificazioni che organizzano queste applicazioni, distinguendo tra AI "meccanica", "cognitiva" e "affettiva" e mostrando come le capacità proprie dell'AI si traducano in vari tipi di benefici (Huang & Rust, 2021). In modo coerente, altri contributi sottolineano che l'AI sta cambiando il marketing e la comunicazione poiché agisce su compiti e decisioni sia operative che strategiche (Davenport et al., 2020).

In questo contesto, quando si osserva dove la ricerca empirica si è maggiormente concentrata, è subito evidente che la maggior parte della letteratura ha trattato l'AI come leva di comunicazione focalizzandosi soprattutto in ambienti digitali, in cui la grande disponibilità di dati e la facile misurabilità degli outcome rende più semplici test e validazioni, lasciando meno esplorate le applicazioni in media non-digitali o tradizionali.

2.2.1 AI come infrastruttura data-driven per pianificare, targettizzare e ottimizzare la comunicazione

Un importante filone di ricerca in materia di AI e comunicazione studia l'AI al centro della comunicazione data-driven. In questo senso, l'AI permette di selezionare il pubblico più adatto, personalizzare i messaggi, ottimizzare la delivery e l'allocazione del budget e misurare la performance in maniera estremamente efficace ed efficiente, affermandosi come uno strumento capace di generare forme avanzate di targeting e adattamento dinamico dei contenuti (Davenport et al., 2020; Huang & Rust, 2021).

Personalizzazione e targeting algoritmico

Molti studi in materia si concentrano sulla personalizzazione, considerata un beneficio chiave dell'AI perché permette facilmente di adeguare contenuti e offerte a preferenze individuali o piccoli segmenti di consumatori. La letteratura evidenzia però che la personalizzazione non produce solo effetti positivi: da un lato aumenta la pertinenza e l'utilità percepita, dall'altro può generare sensazioni di intrusività e preoccupazioni legate alla privacy, creando il cosiddetto "personalization paradox" (Aguirre et al., 2015).

Altri lavori focalizzati sulla pubblicità online personalizzata mostrano che l'efficacia della personalizzazione dipende da condizioni contestuali, suggerendo che non esiste un effetto univoco ma una struttura più complessa influenzata dalla rilevanza situazionale e dalle percezioni dei consumatori (Bleier & Eisenbeiss, 2015).

Computational advertising: automazione della delivery e ottimizzazione in tempo reale

Nella letteratura su advertising e media planning, la nozione di computational advertising descrive un ecosistema in cui le decisioni di acquisto, la collocazione e l'ottimizzazione degli annunci sono

gestite da sistemi algoritmici, mettendo in luce come l'advertising contemporaneo stia diventando sempre più computazionale.

In questo contesto, Huh e Malthouse (2020) sottolineano come gli algoritmi stiano ridefinendo il campo dell'advertising, chiarendo la necessità di avere nuove agende di ricerca su efficacia, trasparenza e governance in questo senso.

AI e customer experience: automazione conversazionale e interazioni brand-consumer

Un filone di ricerca sulle applicazioni dell'Intelligenza Artificiale riguarda l'uso dell'AI come strumento di interfaccia comunicativa attraverso chatbot e agenti conversazionali. In questo contesto, la letteratura mostra che le caratteristiche antropomorfe e gli elementi comunicativi dell'agente artificiale che si interfaccia con gli utenti influenzano le percezioni degli utenti stessi e la valutazione dell'azienda, incidendo sui risultati attitudinali (Araujo, 2018). Questo punto si inserisce nella più ampia prospettiva esperienziale proposta da Puntoni et al. (2021), che colloca l'AI all'interno del customer journey come un insieme di esperienze che determinano come il interpreta l'intervento dell'AI.

Reazioni del consumatore all'AI: accettazione, resistenza e attribuzione di agency

La letteratura evidenzia che l'utilizzo di AI genera reazioni psicologiche nella mente dei consumatori, che in alcuni casi mostrano algorithm appreciation, in altri algorithm aversion. Questo è particolarmente rilevante perché implica che l'utilizzo di AI nella comunicazione può influenzare i comportamenti degli utenti, moderando l'efficacia dei messaggi e influenzandone la credibilità. Alcuni studi sperimentali dimostrano che la distinzione tra algorithm appreciation e algorithm aversion da parte dei consumatori può essere riscontrata in vari contesti di marketing e advertising, con le indagini che si sono focalizzate, soprattutto, sullo studio della risposta dei consumatori rispetto a media e scenari digitali.

2.2.2 AI come motore creativo e strumento di produzione e trasformazione dei contenuti

Un secondo grande filone di ricerca in materia, cresciuto rapidamente con l'emergere della Generative AI, si focalizza sull'AI come strumento capace di produrre messaggi e contenuti pubblicitari. In

questo senso, la letteratura sottolinea che la GenAI influenza la comunicazione di marketing in tre modi principalmente: permette una produzione di contenuti più rapida e scalabile, facilita la personalizzazione e l'adattamento dei contenuti e influenza le percezioni dei consumatori, enfatizzando sia benefici che rischi (Grewal et al., 2025).

Da un punto di vista manageriale e teorico, la GenAI è interpretata in letteratura come uno strumento in grado di ristrutturare i processi creativi tipici del marketing, con effetti di efficienza e potenzialmente anche di efficacia (Grewal et al., 2025). La ricerca sui contenuti AI-generated mostra, però, che gli effetti dell'AI generativa non dipendono solo dal contenuto in sé, ma anche dalla conoscenza della sua origine, ovvero dalla consapevolezza da parte dei consumatori che quel contenuto sia stato creato da AI. In ambito pubblicitario, la disclosure dell'uso di AI può modificare i processi cognitivi degli utenti, influenzandone la fiducia, la valutazione dell'annuncio e attivando eventualmente reazioni difensive. Ad esempio, alcuni studi indicano che dichiarare esplicitamente che il contenuto sia stato generato da AI può incidere sulla valutazione dell'ad e sulle risposte dei consumatori, suggerendo che la disclosure agisce come segnale interpretativo che cambia l'intera elaborazione del messaggio (Baek et al., 2024).

Un risultato che sta emergendo molto spesso dalla letteratura recente è che l'uso di GenAI per creare contenuti può generare benefici di efficienza, ma anche costi. In particolare, quando i consumatori scoprono o sospettano che il contenuto sia stato creato da AI, possono pensare che sia stato coinvolto un minore sforzo umano, che l'annuncio sia meno genuino o che dietro al messaggio ci sia un intento più strumentale, portando ad una diminuzione dell'autenticità percepita e, in alcuni casi, a peggiori valutazioni del brand. Evidenze sperimentali, inoltre, confermano che i contenuti creati con GenAI, in contesti social e digitali, possono diminuire l'autenticità percepita del brand, segnalando un potenziale trade-off tra questo meccanismo e la maggiore produttività portata dall'AI (Brüns & Meißner, 2024).

2.2.3. Una ricerca sbilanciata verso i media digitali

La maggior parte della letteratura esistente sull'utilizzo di AI in contesti di marketing e comunicazione è focalizzata su ambienti e media digitali. Le review della letteratura mostrano che le evidenze in materia riguardano soprattutto aree come digital marketing integrato, contenuti digitali, customer interaction e market research (Chintalapati & Pandey, 2022) e, più in generale, mostrano

che molti contributi sono relativi a social media, e-commerce e automazione di servizi e recensioni (Mariani et al., 2022). Allo stesso modo, le agende di ricerca sul computational advertising mostrano che i temi più sviluppati, come targeting, delivery e misurazione, si focalizzano soprattutto in ambienti digitali (Huh & Malthouse, 2020).

Questa predominanza dei media digitali può essere spiegata dal fatto che, quando la comunicazione avviene in contesti digitali, per l'AI è più facile ottenere informazioni, creare personalizzazioni e generare ottimizzazioni rapidamente, rendendo più semplice la conduzione di studi sperimentali.

2.2.4. Oltre il digitale: dove la letteratura esiste ma è più frammentata

La maggior parte delle evidenze esistenti su AI e comunicazione di marketing è focalizzata su contesti prevalentemente digitali, mentre le applicazioni dell'AI nei media fisici e ambientali risultano essere meno presenti nella letteratura. Esistono alcune ricerche su digital signage e comunicazione in-store (che può essere visto come un ponte tra fisico e digitale), che indagano gli atteggiamenti verso la pubblicità su schermi e le condizioni di valutazione del messaggio nel punto vendita (Lee & Cho, 2019; van de Sanden et al., 2020), ma questi studi si concentrano su variabili classiche dell'advertising e non sui costrutti specifici dell'AI.

Tra i media ambientali poco studiati in questo senso, un media particolarmente rilevante è l'Out-of-Home (OOH) e, in particolare, la sua evoluzione digitale, il Digital Out-of-Home (DOOH).

Esistono alcuni studi in materia, ma rispetto alla quantità delle ricerche sull'AI in contesti digitali, la letteratura che tratta direttamente l'Intelligenza Artificiale in contesti di OOH e DOOH rimane più frammentata e meno sviluppata.

In sintesi, la letteratura sulle applicazioni dell'AI nella comunicazione di marketing risulta essere molto concentrata sui media digitali, il che rende particolarmente rilevante estendere l'analisi a contesti meno studiati come l'OOH e il DOOH.

Da qui emerge quasi in modo naturale un primo gap di ricerca, che la presente tesi si propone di coprire: mentre la letteratura su AI e comunicazione di marketing è ricca nei media digitali e inizia a consolidarsi sui temi di GenAI, l'applicazione dell'AI nella comunicazione OOH e DOOH, un contesto fisico, pubblico e caratterizzato da dinamiche diverse, rimane meno esplorata e poco integrata con i framework prevalenti nella ricerca.

2.3 OOH e DOOH nella letteratura: caratteristiche del mezzo, determinanti di efficacia e implicazioni teoriche

2.3.1 OOH come mezzo pubblicitario: attenzione, memoria e contesto

L'Out-of-Home (OOH) è definito come un mezzo ad ampia copertura e alta frequenza potenziale, inserito in un contesto urbano. Dal momento che l'esposizione agli annunci OOH avviene durante altre attività, gli utenti assistono al messaggio in condizioni di attenzione limitata, bassa elaborazione e vincoli temporali e operano con risorse cognitive limitate e in presenza di molteplici altri stimoli ambientali. Questo implica che la performance del mezzo non dipende soltanto dalla qualità degli annunci, ma dalla capacità degli stessi di guadagnare una parte sufficiente dell'attenzione degli utenti per essere notato e ricordato.

In questo senso, la letteratura sull'attenzione pubblicitaria fornisce alcuni elementi che sono molto importanti per comprendere cosa determina l'efficacia OOH. Alcuni studi mostrano che l'attenzione verso l'annuncio è influenzata da fattori come dimensione degli elementi, complessità e struttura visiva (Pieters & Wedel, 2004) e che una comunicazione OOH efficace deve massimizzare fluidità e salienza dei contenuti e ridurre il carico cognitivo necessario alla comprensione dell'annuncio, aumentando così la probabilità di memorizzarlo e ricordarlo. Inoltre, uno studio empirico sull'outdoor advertising (Wilson, Baack e Till, 2015) evidenzia che la creatività e la visibilità sono prerequisiti necessari per la memoria di marca, mostrando che la creatività da sola non è sufficiente se l'annuncio non risalta abbastanza da catturare attenzione.

Un aspetto centrale dell'OOH, inoltre, è la relazione tra esposizione ripetuta e memoria. Anche in assenza di un'elaborazione profonda, la ripetizione dell'annuncio può aumentare familiarità e riconoscimento (Zajonc, 1968) ed è quindi importante mostrare ripetutamente gli annunci al pubblico lungo percorsi ricorrenti. Nonostante questo, però, la ripetizione non implica necessariamente efficacia, poiché in un ambiente pieno di stimoli, un annuncio può essere visto molte volte senza essere realmente notato e ricordato e può addirittura generare irritazione se percepito come intrusivo. Inoltre, nell'OOH l'esposizione agli annunci avviene in spazi pubblici ed è involontaria e difficile da evitare. Questo dà molta rilevanza a costrutti come intrusività percepita e reactance, che nella letteratura sono associati a situazioni in cui il consumatore percepisce un qualche tipo di minaccia alla propria libertà personale (Brehm, 1966; Edwards et al., 2002), e sottolinea che nell'OOH è essenziale riuscire a limitare il rischio di generare fastidio o di risultare troppo invasivi.

2.3.2 La frammentazione della letteratura OOH

La letteratura accademica sull'OOH è tendenzialmente meno sviluppata rispetto a quella su altri media come TV, stampa o advertising digitale, risultando tendenzialmente più frammentata. Questa frammentazione implica che molte evidenze, che potrebbero essere rilevanti, non sempre sono integrate in un framework generale di advertising e comunicazione. Inoltre, i cambiamenti dovuti al passaggio da OOH a DOOH implicano un cambiamento anche dei processi psicologici e delle metriche di efficacia collegate al mezzo e la letteratura esistente, basata solo su OOH statico, potrebbe non essere sufficiente per spiegarle. In questa ottica, una review della letteratura (Wilson, 2023) propone una research agenda per ampliare la letteratura in materia di OOH e sottolinea la necessità di approfondire le nuove possibilità legate alla digitalizzazione.

2.3.3 DOOH: natura del mezzo, evidenze di letteratura e opportunità di ricerca

Inquadramento teorico del DOOH

Il Digital Out-of-Home (DOOH) rappresenta l'evoluzione dell'OOH verso forme più dinamiche e digitali, che implicano l'utilizzo di schermi digitali per mostrare gli annunci, una rotazione creativa dei contenuti più rapida e la possibilità di adattare i contenuti a variabili contestuali, come ora del giorno, meteo, eventi, traffico, densità. In questo senso potremmo considerare il DOOH come un mezzo ibrido: rimane un media ambientale caratterizzato dallo spazio pubblico e dall'esposizione rapida e incidentale tipica dell'OOH, ma allo stesso tempo possiede caratteristiche tipiche dei media digitali, come dinamicità, aggiornamento in tempo reale e potenziale data-driven. Il DOOH può quindi approfittare dei benefici attribuiti ai nuovi media, come maggiore pertinenza e possibilità di ottimizzazione, mantenendo comunque le sue proprietà di ampia copertura e presenza fisica.

Questa natura del DOOH implica che la digitalizzazione non cambia solo come vengono prodotti o mostrati i contenuti, ma potrebbe influenzare anche i comportamenti e le percezioni dei consumatori. Un annuncio DOOH pertinente rispetto al contesto può essere percepito come più utile e rilevante, perché appare adatto al momento e al luogo, ma può anche essere percepito come intrusivo, perché il consumatore potrebbe interpretare l'adattamento al contesto come una sorveglianza. Il DOOH rientra

quindi in un trade-off che appare spesso quando si parla di personalizzazione: una maggiore rilevanza può aumentare l'efficacia, ma le domande su come e perché il messaggio sia così rilevante possono generare resistenze e percezioni negative.

Evidenze e filoni della letteratura su DOOH

La ricerca accademica sul DOOH è tendenzialmente frammentata, ma tratta comunque alcuni filoni di ricerca particolarmente rilevanti per comprendere al meglio il mezzo.

Un primo filone tratta il collegamento tra l'efficacia del DOOH e le sue caratteristiche di attenzione limitata ed esposizione rapida. Alcuni studi mostrano che, in contesti di esposizione rapida e urbana, il fatto che i consumatori sono sottoposti a tanti stimoli potrebbe ridurre la loro capacità di riconoscere e ricordare il brand visto in un annuncio e suggeriscono che per migliorare questo aspetto è importante creare contenuti accattivanti e facilmente leggibili (van Meurs & Aristoff, 2009); altri studi sul billboard advertising mostrano che, in ambienti urbani saturi di stimoli, un annuncio per essere efficace deve catturare l'attenzione ed essere particolarmente visibile (Wilson et al., 2015).

Un secondo filone riguarda i benefici sperimentali che il DOOH può ottenere grazie al suo essere più digitalizzato e afferma che alcuni elementi, come l'utilizzo di schermi digitali e le rotazioni frequenti, permettono al DOOH di essere più vicino alle logiche sperimentali e alla possibilità di condurre esperimenti di campo. A favore di questo punto esistono alcuni studi che mostrando come il DOOH sia più adatto per misurare effetti causali rispetto a molte forme di OOH tradizionale (Lee, 2025). In generale, sta emergendo un filone di ricerca sul DOOH quantitativo e ricco di dati, che sta studiando le proprietà digitali del DOOH per misurarne l'efficacia.

Un terzo filone che si può identificare non riguarda direttamente il DOOH, ma è comunque molto rilevante per comprendere al meglio le dinamiche del mezzo. Si tratta della letteratura sul digital signage in contesti fisici, che studia il valore percepito e gli atteggiamenti dei consumatori verso la pubblicità su schermi digitali in-store. Uno studio sul targeting emozionale, effettuato attraverso strumenti di face-reading, mostra che la personalizzazione dei contenuti e dei messaggi aumenta le intenzioni di acquisto e le valutazioni del prodotto, ma che la disclosure su come è stato effettuato il targeting elimina gli effetti positivi perché fa aumentare la percezione che dietro al messaggio ci sia un intento manipolativo (Garaus et al., 2021). Questi risultati ci indicano che la trasparenza sull'utilizzo di strumenti tecnologici può influenzare pesantemente l'interpretazione del messaggio e la sua l'efficacia.

AI e DOOH: un'applicazione ancora poco studiata

Come già detto, la maggior parte degli studi empirici sull'utilizzo di AI nella comunicazione di marketing è concentrata sugli ambienti puramente digitali, mentre l'applicazione dell'AI in ambienti fisici digitalizzati, come appunto il DOOH, è meno presente nella letteratura.

Esistono molti contributi sul DOOH come mezzo e molti contributi sull'AI nella comunicazione di marketing, ma mancano ancora degli studi che integrano questi due filoni e che testano in modo causale come gli elementi tipici dell'AI operano nel contesto DOOH. Anche la research agenda sull'OOH proposta da Wilson (2023) è d'accordo con questo punto e, infatti, parla della necessità di integrare la letteratura su OOH e DOOH rispetto alle recenti evoluzioni tecnologiche.

AI applicata al DOOH: direzioni di ricerca basate sulla letteratura

Dalla letteratura esistente sul DOOH e sull'AI, emergono alcune direzioni di ricerca che potrebbero essere efficaci per integrare i due filoni. Tra queste, ci possiamo soffermare su tre in particolare.

Una prima direzione si basa sull'utilizzo dell'AI per la creazione di contenuti DOOH e si concentra sul concetto di origine creativa. Alcune evidenze in ambito social e digitale indicano che l'uso di GenAI per la creazione di contenuti può ridurre le percezioni di autenticità da parte dei consumatori, con effetti potenzialmente negativi sulle loro valutazioni e la loro relazione con il brand (Brüns & Meißner, 2024). Questo contributo suggerisce che l'origine creativa AI-generated dei contenuti può essere percepita in due modi: come innovazione, generando conseguenze positive, oppure come artificialità e non autenticità, portando invece a risvolti negativi. Questa doppia percezione ha delle conseguenze su importanti outcome che misurano l'efficacia pubblicitaria, come l'atteggiamento verso l'annuncio (Aad) e la brand memory. (Wilson et al., 2015; van Meurs & Aristoff, 2009).

Una seconda direzione di ricerca tratta il ruolo della disclosure dell'uso di AI nel DOOH e le sue potenziali conseguenze. La letteratura mostra che la disclosure dell'utilizzo di strumenti tecnologici può influenzare l'interpretazione del messaggio e ridurre l'efficacia quando attiva sospetti o percezione che il messaggio abbia un intento manipolativo (Garaus et al., 2021). Altri studi indicano che la disclosure dell'uso di AI può peggiorare gli atteggiamenti verso l'annuncio in alcune condizioni (Grigsby et al., 2025), suggerendo un trade-off secondo il quale l'etichetta di disclosure può essere interpretata come un elemento di trasparenza oppure come un attivatore di diffidenza. Applicare

questi risultati al DOOH potrebbe essere particolarmente rilevante perché lo spazio pubblico amplifica l'importanza dei temi di etica e controllo.

Una terza ed ultima direzione si origina a partire dalla capacità dell'AI di personalizzare e contestualizzare i messaggi. Il DOOH permette la creazione di contenuti adattivi e l'utilizzo di AI potrebbe essere sfruttato per amplificare questa capacità, con l'ottimizzazione in tempo reale dei contenuti e l'adattamento del messaggio al contesto. In teoria, questo dovrebbe aumentare utilità e pertinenza percepita, dal momento che il consumatore si trova davanti ad un messaggio coerente rispetto alla situazione in cui si trova, ma potrebbe anche aumentare l'intrusività percepita se l'adattamento viene interpretato come un segnale di sorveglianza. La letteratura suggerisce che la personalizzazione del messaggio può funzionare, ma dire esplicitamente come avviene questa personalizzazione può invertirne l'effetto (Garaus et al., 2021), portando ad un trade-off tra beneficio e costo.

Da queste direzioni di ricerca possiamo capire il ruolo che l'AI potrebbe assumere in relazione al DOOH e possiamo identificare tre elementi che potrebbero essere interessanti da studiare in questo contesto: l'origine creativa, la disclosure e la pertinenza contestuale. Queste tre leve sono state studiate soprattutto in ambienti digitali, mentre nel DOOH gli studi sono molto limitati e spesso indiretti.

Da qui, possiamo identificare un duplice gap. In primo luogo, nella letteratura mancano dei contributi che studiano gli effetti dell'origine creativa AI-generated dei contenuti e della disclosure dell'uso di AI nel DOOH. In secondo luogo, mancano evidenze sugli effetti della pertinenza contestuale del messaggio nel DOOH, con la relativa analisi del trade-off tra utilità e rilevanza da una parte e intrusività dall'altra.

2.4 AI e advertising: origine creativa e disclosure come segnali persuasivi

2.4.1 Origine creativa dei contenuti: "AI-generated" vs "human-created"

L'origine creativa di un contenuto è "chi" o "cosa" ha creato quel contenuto ed è un elemento che può influenzare le valutazioni dei consumatori e l'interpretazione del messaggio.

La creazione di messaggi e contenuti pubblicitari e la creatività in generale sono elementi che l'opinione pubblica attribuisce agli esseri umani e quindi, quando un annuncio è percepito come AI-generated e la sua origine creativa è attribuita all'AI, i consumatori potrebbero pensare che ci sia stato un minore contributo umano nella realizzazione dell'annuncio e questo potrebbe avere degli effetti negativi sulla valutazione dell'annuncio stesso. Ci sono anche delle evidenze empiriche che supportano questo punto, mostrando che l'uso di GenAI per creare contenuti può ridurre l'autenticità percepita del brand e può essere visto come un segnale di minore genuinità (Brüns & Meißner, 2024). Tuttavia, la valutazione degli annunci AI-generated non è sempre negativa, ma dipende dal tipo di messaggio che si crea e da cosa le persone pensano che l'AI dovrebbe o non dovrebbe fare. Chen et al. (2024) mostrano che i consumatori reagiscono meglio ad annunci AI-generated quando il messaggio è orientato a elementi tecnici come performance, efficacia e risultati, mentre preferiscono annunci human-created quando il messaggio è creativo o ha un obiettivo comunicativo. L'effetto dell'origine creativa AI, quindi, dipende in gran parte dal fatto che il contenuto creato sia allineato o meno alle cose che le persone pensano che l'AI possa fare portando ad un dualismo tra la percezione dell'AI come moderna, innovativa e competente e la percezione dell'AI come meno umana e meno autentica.

2.4.2 Disclosure dell'uso dell'AI: trasparenza, credibilità e effetti negativi

La disclosure, quando usata all'interno di messaggi, contenuti o annunci pubblicitari in contesti di marketing e advertising, è un elemento che rende esplicite alcune informazioni riguardo il contenuto o l'annuncio. La letteratura in materia indica che la disclosure non è, però, solo un elemento informativo, ma è un elemento che può cambiare il modo in cui gli utenti e i consumatori percepiscono il messaggio. Una parte della letteratura mostra che la presenza di disclosure di solito aumenta la percezione che l'annuncio abbia uno scopo persuasivo, influenzando in modo negativo gli atteggiamenti dei consumatori e le loro valutazioni nei confronti del contenuto e del brand.

La disclosure dell'uso di AI informa i consumatori sul fatto che il contenuto sia stato creato da uno strumento artificiale, il che ha delle conseguenze negative sull'affidabilità e l'autenticità percepite. In particolare, la letteratura mostra che la disclosure dell'uso di AI per la creazione di contenuti può generare valutazioni meno favorevoli dell'annuncio (Baek, Kim e Kim, 2024), può diminuire la fiducia verso il brand e il messaggio e può peggiorare l'atteggiamento dei consumatori verso l'ad (Grigsby, Michelsen e Zamudio, 2025).

Tuttavia, la disclosure non produce necessariamente solo effetti negativi e infatti la letteratura evidenzia che gli effetti della disclosure non sono univoci, ma dipendono dal tipo di disclosure e dal contesto. Henestrosa e Kimmerle (2025) dimostrano che disclosure più informative che includono limiti e punti di forza possono anche aumentare la credibilità percepita, perché in quel caso la disclosure è vista come un segnale di trasparenza e correttezza. Inoltre, secondo la cosiddetta machine heuristic, le persone attribuiscono alle macchine imparzialità e oggettività e le percepiscono meno soggette a bias umani e più credibili quando il compito che svolgono è oggettivo o tecnico (Sundar & Kim, 2019; Yang et al., 2024). Se applichiamo questa cosa all'AI nell'advertising, possiamo ipotizzare che la disclosure dell'uso di AI potrebbe aumentare la credibilità percepita quando il claim dell'annuncio è di tipo tecnico o oggettivo.

2.4.3 Origine creativa, disclosure e implicazioni per il DOOH

L'origine creativa e la disclosure sono elementi che influenzano le reazioni dei consumatori e le loro interpretazioni di messaggi e annunci di marketing. In un contesto come il DOOH questi meccanismi sono ancora più forti perché le valutazioni dei consumatori, a causa dell'esposizione breve e dell'attenzione limitata, sono basate sulla prima impressione e elementi come origine creativa AI e disclosure dell'uso di AI potrebbero condizionarla molto, aumentando il peso di tali elementi e rendendo più probabili i meccanismi sottostanti. Per questo, per capire bene l'intersezione tra AI e DOOH è importante studiare gli effetti dell'origine creativa e della disclosure.

2.5 Cornici teoriche per spiegare gli effetti nel DOOH: persuasione, segnali, conoscenza della persuasione e reactance

Per spiegare bene gli effetti dell'AI nel DOOH, è necessario sintetizzare alcuni punti teorici che potrebbero essere utili per spiegare la relazione. Per fare questo, la presente tesi utilizza quattro cornici concettuali, l'Elaboration Likelihood Model (ELM), la Signaling Theory, il Persuasion Knowledge Model (PKM) e la Reactance Theory. Queste, se integrate tra di loro, possono indicare quali meccanismi spiegano la relazione tra AI e DOOH e, quindi, quali sono gli elementi che devono essere presi in considerazione per studiare la relazione.

2.5.1 Elaboration Likelihood Model: Dual-process e vincoli di elaborazione in condizioni di “fast exposure”

L'Elaboration Likelihood Model (ELM) distingue tra una route centrale (elaborazione profonda) e una route periferica (uso di scorciatoie cognitive), sottolineando che la direzione e la qualità dell'elaborazione dipendono da motivazione e capacità (Petty & Cacioppo, 1986). In un contesto DOOH, in cui l'esposizione è rapida e ci sono molte distrazioni, il consumatore ha risorse cognitive limitate per elaborare il messaggio e la route periferica gioca un ruolo importante, perché l'utente valuta rapidamente l'annuncio, senza analizzarne i dettagli in profondità.

Questo implica che la forma dell'annuncio e i segnali che contiene possono influenzare molto gli atteggiamenti degli utenti. In questo senso, l'origine creativa e la disclosure diventano particolarmente rilevanti, perché forniscono un'informazione semplice e immediata da elaborare, che può quindi influenzare le reazioni degli utenti.

2.5.2 Signaling theory: disclosure e origine creativa come segnali strategici

La Signaling Theory, spiega che, in condizioni di informazione asimmetrica, vengano usati dei segnali per comunicare le informazioni non direttamente osservabili. Kirmani e Rao (2000) trasportano la teoria nel contesto di marketing e mostrano che alcune variabili del marketing mix possono essere segnali che indicano qualità non osservabili.

Applicando questa teoria al contesto d'interesse, l'origine creativa e la disclosure possono essere interpretate come segnali che comunicano aspetti non osservabili del messaggio, il processo produttivo, la creatività e la trasparenza. Sono, però, dei segnali che possono essere interpretati in maniera ambivalenti, perché, da una parte, l'origine creativa AI può segnalare innovazione, modernità e capacità tecnologica, ma anche standardizzazione e minore contributo umano e, dall'altra, la disclosure dell'uso di AI può segnalare trasparenza e responsabilità, ma può anche generare diffidenza.

La Signaling Theory è quindi utile per affermare che l'origine creativa e la disclosure sono segnali importanti e anche per capire che l'efficacia dipende dall'interpretazione di questi segnali.

2.5.3 Persuasion Knowledge Model: riconoscimento dell'intento persuasivo

Il Persuasion Knowledge Model (PKM) dice che, quando i consumatori riconoscono un intento persuasivo, possono attivare meccanismi che modificano le valutazioni dell'annuncio (Friestad & Wright, 1994). In questo senso, la disclosure "creato con IA" può aumentare sospetti e domande, rendendo più probabile l'attivazione della persuasion knowledge e, quindi, un atteggiamento meno favorevole.

Nel DOOH, anche se l'utente non ha molto tempo di elaborazione, un segnale semplice come la disclosure può essere sufficiente per attivare persuasion knowledge e quindi influenzare le reazioni e gli atteggiamenti dei consumatori.

2.5.4 Reactance theory: minaccia alla libertà, intrusività e backlash

La Reactance Theory definisce la reactance come uno stato motivazionale negativo che emerge quando una libertà viene percepita come minacciata o ridotta e può produrre resistenza e rifiuto (Brehm, 1966). Nei contesti di advertising, la reactance è spesso associata a forme di esposizione percepite come forzate, invasive o intrusive (Edwards, Li e Lee, 2002).

Nel DOOH, il consumatore non sceglie di vedere all'annuncio e l'esposizione è difficile da evitare e, per questo, è più probabile che un messaggio venga percepito come invadente e intrusivo, attivando la reactance. Inoltre, l'AI è in grado di generare contenuti con alta pertinenza contestuale, che può essere interpretata come un segnale di tracciamento e di sorveglianza e può quindi amplificare le percezioni di intrusività e la reactance.

2.6 Elementi chiave del modello: definizioni, rilevanza e stato della ricerca

Nel corso della revisione sviluppata nel presente capitolo, sono emersi alcuni elementi che si inseriscono direttamente nella relazione tra AI e DOOH e che sono fondamentali per studiare i meccanismi che spiegano gli effetti dell'AI nel DOOH.

2.6.1 Attitude toward the Ad (Aad)

L'atteggiamento verso l'annuncio (Attitude toward the Ad, Aad) è una misura chiave dell'efficacia pubblicitaria che riesce a catturare la valutazione complessiva dell'annuncio. Se si parla di contenuti creati da AI, l'Aad è una misura particolarmente rilevante perché riesce a catturare sia le valutazioni più esplicite, come stile, coerenza e qualità percepita, sia le valutazioni cognitive e quindi potrebbe essere adatto per misurare i benefici e i costi che possono derivare da un annuncio AI-generated.

La maggior parte degli studi sull'Aad è stata svolta in contesti tradizionali come televisione e stampa e in contesti digitali e mancano, invece, evidenze dirette che studino l'Aad in contesti di ambientali DOOH, soprattutto in relazioni ad annunci AI-generated.

2.6.2 Memoria di brand: brand recognition

Nel DOOH l'esposizione agli annunci avviene in concorrenza con altri elementi e con attenzione limitata e, per questo, l'obiettivo principale della comunicazione non è la conversione immediata (quasi impossibile vista la struttura del DOOH), ma l'encoding in memoria di elementi importanti presenti nell'annuncio, come nome brand, claim e visual, in modo che poi il consumatore possa riconoscere il brand. La brand recognition è, quindi, un risultato molto importante, ma allo stesso tempo molto difficile da ottenere, a causa dell'attenzione limitata e dell'esposizione rapida.

La memoria di brand dipende dal grado di attenzione e dalle caratteristiche dell'annuncio e la letteratura mostra che un annuncio DOOH ha più probabilità di creare brand recognition se ha un visual semplice e contiene poche informazioni, ma chiare (van Meurs & Aristoff, 2009).

L'utilizzo di AI permette di creare annunci DOOH più dinamici e pertinenti con il contesto e questo potrebbe facilitare l'encoding in memoria e la brand recognition; dall'altra parte, però, un messaggio troppo mirato potrebbe risultare intrusivo, generare resistenza e, quindi, ridurre la memoria, portando a degli effetti non univoci.

Esistono varie evidenze in letteratura sulla memoria di brand in contesti di DOOH, ma ancora non sono stati studiati del tutto gli effetti che gli elementi tipici dell'AI (origine creativa, disclosure e pertinenza contestuale) potrebbero avere sull'encoding in memoria e la brand recognition.

2.6.3 Autenticità, fiducia e credibilità

Nel marketing, l'autenticità riguarda le valutazioni che il brand sia vero e sincero e la percezione di autenticità influenza molto gli esiti dell'efficacia pubblicitaria. Questo può essere ampliato quando si studiano gli effetti dell'AI, perché l'utilizzo di AI per la creazione di contenuti e messaggi, come già detto, può essere percepito come una riduzione del contributo umano e questo può avere un impatto negativo sull'autenticità percepita e, quindi, sull'efficacia. L'autenticità percepita è poco studiata in relazione al DOOH e, ancora meno, se unito all'AI; tuttavia, studiare l'autenticità percepita in questo senso potrebbe risultare particolarmente rilevante, perché l'origine creativa AI del contenuto e la presenza di disclosure possono cambiare la percezione del contenuto e l'effetto dell'autenticità potrebbe essere importante.

La credibilità percepita e la fiducia, invece, riguardano le percezioni che un messaggio o fonte siano affidabili e competenti. Quando entra in gioco l'AI, in particolar modo se il suo utilizzo è dichiarato, nei consumatori possono emergere dubbi e domande riguardo il processo di creazione e l'intento del contenuto e questo può impattare negativamente la credibilità del contenuto e la fiducia verso il brand, che diventano meccanismi che possono influenzare l'efficacia.

2.6.4 Persuasion knowledge, reactance e resistenza persuasiva

La persuasion knowledge è la consapevolezza o l'attribuzione da parte dei consumatori di intenti persuasivi di un contenuto e può portare a degli atteggiamenti di resistenza per contrastare questi intenti di influenza (Friestad & Wright, 1994). La reactance, invece, è la risposta a minacce percepite alla propria libertà di scelta, portando anche qui a reazioni difensive che possono avere effetti negativi sull'efficacia dell'annuncio (Brehm, 1966).

L'unione tra AI e DOOH è un contesto in cui i meccanismi di reactance e persuasion knowledge possono avvenire molto facilmente, per due motivi. In primis, l'AI permette di creare messaggi pubblicitari pertinenti con il contesto e questa pertinenza potrebbe essere interpretata come sorveglianza, aumentando la probabilità di reactance e persuasion knowledge; inoltre, nel DOOH l'utente ha poco controllo e spesso non può evitare l'esposizione all'annuncio e questo potrebbe aumentare la probabilità di attivare reazioni difensive.

2.6.5 Pertinenza e utilità

Nella comunicazione di marketing, la relevance (pertinenza) indica quanto il messaggio sia appropriato rispetto a bisogni e obiettivi del consumatore; l'utility (utilità), invece, riguarda il valore percepito che il messaggio offre, valore che può essere di vario tipo, informativo, funzionale o esperienziale. Il modello dell'advertising value studia alcuni elementi che influenzano l'utilità percepita di un annuncio, collegando fattori come informazione e intrattenimento ad un valore percepito maggiore e fattori come fastidio o irritamento ad un valore minore (Ducoffe, 1996). Inoltre, la letteratura mostra che la congruenza tra messaggio e contesto è un fattore che influisce positivamente sul valore, aumentando pertinenza e utilità (Hühn et al., 2017).

Nel DOOH, una delle funzioni principali dell'AI è la sua capacità di adattare i contenuti al contesto e questo potrebbe aumentare l'efficacia del mezzo grazie all'aumento di relevance e utility.

Nel complesso, nella letteratura sono presenti molte evidenze che studiano questi elementi di interesse; manca, però, un'integrazione di questi elementi nel contesto dell'applicazione dell'AI al DOOH.

2.7 Dalla letteratura al gap

Dalla revisione fatta nel corso del capitolo emerge il gap di ricerca che la presente tesi intende colmare: nella letteratura esistente, mancano degli studi empirici che studino sistematicamente l'applicazione e l'utilizzo dell'Intelligenza Artificiale nel contesto del DOOH e che testino quali condizioni rendono queste applicazioni efficaci o meno. In particolare, il presente lavoro si propone di colmare un gap specifico: se e quando le caratteristiche AI-based del DOOH, come origine creativa, disclosure e pertinenza context-aware, rendono il DOOH più efficace o, al contrario, potenzialmente controproducente in termini di efficacia pubblicitaria e attraverso quali meccanismi psicologici.

In questo contesto, il DOOH permette di studiare cosa succede quando segnali tipici dell'AI advertising, studiati soprattutto nel digitale, vengono trasferiti in un media tradizionale e ambientale, con caratteristiche di esposizione e attenzione molto diverse.

2.8 Dal gap alle domande di ricerca

A partire dal gap identificato attraverso la revisione, le domande di ricerca del modello sono così strutturate:

RQ1 – Efficacia della creatività AI nel DOOH

La creatività DOOH generata con AI è più (o meno) efficace rispetto a una creatività DOOH generata da soggetti umani in termini di: a) attitude toward the ad (Aad) e b) brand recognition?

RQ2 – Ruolo della disclosure “Creato con AI”

La disclosure dell’uso di AI attenua o amplifica l’efficacia della creatività AI nel DOOH (in termini di Aad e brand recognition)? E attraverso quali meccanismi psicologici?

RQ3 – Pertinenza context-aware nel DOOH e interazione con la disclosure

Nel DOOH con creatività AI, l’alta pertinenza context-aware aumenta l’efficacia dell’annuncio (in termini di Aad e brand recognition) e la percezione di utilità e rilevanza? La presenza di disclosure influenza questo effetto?

Nel loro insieme, le tre RQ permettono di colmare il gap precedentemente identificato e di testare nel DOOH i principali costrutti con cui la letteratura ha finora studiato l’AI in advertising e comunicazione e verificare sia il loro effetto sull’efficacia del mezzo, sia i meccanismi che spiegano l’effetto.

2.9 Sviluppo delle ipotesi di ricerca

La presente ricerca sviluppa un modello che studia gli effetti dell’origine creativa (AI-generated vs human-generated) e della pertinenza context-aware (alta vs bassa) sull’efficacia del DOOH, analizzando gli effetti specifici quando l’uso dell’AI non è esplicitato e quando l’uso dell’AI viene dichiarato attraverso la disclosure.

Per testare questi effetti, la ricerca si avvale di un set di nove ipotesi, che formalizzano ciò che ci si può attendere riguardo agli effetti diretti sugli outcome di efficacia pubblicitaria (Aad e brand memory) e sui meccanismi psicologici che li spiegano.

H1 – Effetto diretto dell'origine creativa sull'efficacia del contenuto

H1a: In assenza di elementi che esplichino l'origine creativa, un contenuto DOOH con origine creativa AI-generated e un contenuto DOOH con origine creativa human-generated non sono significativamente differenti in termini di Aad.

H1b: In assenza di elementi che esplichino l'origine creativa, un contenuto DOOH con origine creativa AI-generated e un contenuto DOOH con origine creativa human-generated non sono significativamente differenti in termini di brand recognition.

Nel DOOH l'elaborazione è rapida e basata su risorse cognitive limitate e gli utenti tendono a basarsi sulle caratteristiche immediatamente percepibili dell'annuncio. Se l'origine creativa non è direttamente esplicitata, è difficile che gli utenti possano accorgersi della differenza ed è, quindi, logico ipotizzare che le due origini creative non abbiano effetti diversi in termini di efficacia.

H2 – Effetto della disclosure sull'origine creativa

H2a: In presenza di disclosure “Creato con AI”, la creatività AI-generated risulta meno efficace in termini di Aad (Aad più basso) rispetto ai casi in cui la disclosure non è presente.

H2b: In presenza di disclosure “Creato con AI”, la creatività AI-generated risulta meno efficace in termini di brand recognition (brand recognition inferiore) rispetto ai casi in cui la disclosure non è presente.

La disclosure rende esplicita l'informazione sull'origine creativa dell'annuncio e sposta l'attenzione valutativa del consumatore dal solo contenuto dell'annuncio al modo in cui è stato prodotto, con possibilità di un impatto negativo in termini di efficacia.

H3 – Effetti della consapevolezza dell'origine AI su autenticità, trust e credibilità

H3: La consapevolezza dell'origine creativa AI-generated di un annuncio riduce: a) autenticità percepita, b) fiducia e c) credibilità percepita rispetto alle condizioni in cui l'origine AI non è nota.

Quando il consumatore sa che il contenuto è generato da AI, può pensare ad una minore genuinità e minore umanità del messaggio, riducendo l'autenticità percepita. Inoltre, l'utilizzo di AI può essere associato a un intento di persuasione, che riduce la fiducia verso il contenuto e la credibilità percepita.

H4 – Effetti di mediazione di autenticità, trust e credibilità

H4: L'effetto negativo della disclosure sull'efficacia dell'annuncio in termini di Aad è mediato da autenticità percepita, fiducia e credibilità: la presenza di disclosure "Creato con AI" comporta una riduzione di autenticità, fiducia e credibilità, che a sua volta comporta una diminuzione di Aad.

H5 – Effetti della consapevolezza dell'origine AI su persuasion knowledge e reactance

H5: La consapevolezza dell'origine creativa AI-generated di un annuncio aumenta: a) persuasion knowledge e b) reactance rispetto alle condizioni in cui l'origine AI non è nota.

H6 – Effetto diretto della pertinenza context-aware sull'efficacia del contenuto

H6a: A parità di creatività AI, un annuncio DOOH con alta pertinenza context-aware genera un contenuto più efficace in termini di Aad rispetto a un annuncio con bassa pertinenza context-aware.

Nel DOOH quando il messaggio è chiaramente coerente con il contesto, viene percepito come più appropriato e significativo, aumentando la probabilità di valutazioni positive dell'annuncio.

H6b: A parità di creatività AI, un annuncio DOOH con alta pertinenza context-aware genera un contenuto più efficace in termini di brand recognition rispetto a un annuncio con bassa pertinenza context-aware.

L'alta pertinenza contestuale aumenta la probabilità che l'annuncio venga notato e che quindi venga codificato in memoria e ricordato.

H7 – Effetto della pertinenza context-aware su utilità e rilevanza

H7. Un annuncio DOOH AI-generated con alta pertinenza context-aware aumenta utilità e rilevanza percepita rispetto ad un annuncio con bassa pertinenza context-aware.

H8 – Effetto di mediazione positiva di utilità e rilevanza

H8a. L'effetto dell'alta pertinenza context-aware su Aad è mediato positivamente da utilità e rilevanza percepita: un'alta pertinenza comporta maggiore utilità e rilevanza che a sua volta migliorano l'Aad.

H8b. L'effetto dell'alta pertinenza context-aware sulla brand recognition è mediato positivamente da utilità e rilevanza percepita: un'alta pertinenza comporta maggiore utilità e rilevanza che a sua volta migliorano la brand recognition.

H9 – Effetto di moderazione della disclosure

H9a. La presenza di disclosure “Creato con AI” attenua l'effetto positivo dell'alta pertinenza context-aware sull'Aad.

H9b. La presenza di disclosure “Creato con AI” attenua l'effetto positivo dell'alta pertinenza context-aware sulla brand recognition.

La disclosure rende esplicito l'utilizzo dell'AI e può portare i consumatori a interpretare l'adattamento contestuale in maniera negativa, riducendo i benefici dell'alta pertinenza in termini di efficienza.

Capitolo III. Metodologia

3.1 Introduzione al capitolo: obiettivo della ricerca e approccio metodologico

Questo capitolo descrive la metodologia utilizzata per condurre la ricerca empirica svolta all'interno della tesi. Si tratta di una ricerca sperimentale di tipo quantitativo, composta da due studi distinti a cui fanno riferimento due esperimenti tra-soggetti (between-subjects).

La scelta del disegno sperimentale è stata coerente con la necessità di ottenere evidenze causali sul ruolo di tre fattori chiave e di stimarne gli effetti causali: origine creativa del contenuto (AI-generated vs human-generated), disclosure dell'utilizzo dell'AI ("Creato con AI") e pertinenza context-aware (alta vs bassa). Questi fattori sono stati testati su outcome di efficacia pubblicitaria coerenti con il contesto DOOH: attitude toward the ad (Aad) e brand memory (operazionalizzata tramite brand recognition). Inoltre, la ricerca ha integrato nello studio alcuni meccanismi psicologici che potrebbero spiegare la relazione tra DOOH e AI: autenticità percepita, fiducia, credibilità, persuasion knowledge e reactance per lo Studio 1; utilità e rilevanza percepita per lo Studio 2.

Il capitolo presenta: il disegno complessivo e il modello di ricerca; i criteri di selezione dei partecipanti e il campionamento; gli strumenti di ricerca (questionario e stimoli); le procedure di raccolta dati; le tecniche di analisi statistica utilizzate; la preparazione del dataset e i criteri di pulizia e codifica; gli aspetti etici e le limitazioni metodologiche riscontrate, così da garantire replicabilità e rigore.

3.2 Disegno di ricerca

3.2.1 Struttura complessiva

La ricerca è stata articolata in due studi quantitativi coerenti tra di loro, focalizzati su meccanismi diversi e progettati per rispondere ad un sottoinsieme delle domande di ricerca e per testare una parte del set di ipotesi.

Lo Studio 1, incentrato su origine creativa del contenuto e disclosure, ha risposto alle prime due domande di ricerca e ha testato: l'efficacia della creatività DOOH AI-generated vs human-generated e la loro equivalenza in assenza di segnali espliciti (H1a e H1b) e l'effetto negativo della disclosure "Creato con AI" sull'efficacia della creatività AI (H2a e H2b), spiegandolo tramite meccanismi psicologici (H3–H5).

Lo Studio 2, incentrato su pertinenza context-aware e disclosure, ha risposto alla terza domanda di ricerca e ha testato: se, a parità di creatività AI nel DOOH, un'alta pertinenza context-aware aumenti l'efficacia pubblicità (H6 a e H6b) e l'utilità e rilevanza percepita (H7–H8) e se la disclosure attenui o meno tali benefici (H9).

L'articolazione della ricerca in due studi così definiti ha permesso di mantenere il disegno statisticamente pulito e leggibile e di isolare i vari costrutti di interesse.

3.2.2 Tipologia dello studio e giustificazione dell'approccio sperimentale

Entrambi gli studi sperimentali che compongono la ricerca sono stati implementati come esperimenti tra-soggetti (between-subjects), con assegnazione casuale (randomizzazione) delle condizioni e misurazione post-esposizione. Ciascun partecipante è stato assegnato casualmente a una condizione sperimentale ed esposto a un solo stimolo DOOH.

La scelta di questa struttura per gli studi è stata motivata da tre ragioni metodologiche:

- a) Validità interna e causalità: la randomizzazione delle condizioni consente di attribuire le differenze negli outcome alle varie manipolazioni (origine, disclosure, pertinenza), riducendo così elementi di confusione.
- b) Compatibilità con la natura del DOOH: un esperimento between-subjects è particolarmente appropriato per il contesto DOOH, dal momento che riesce a riprodurre meglio le caratteristiche dell'esposizione a un annuncio DOOH (esposizione rapida e spesso unica) rispetto a un within-subjects.
- c) Coerenza con le RQ e le ipotesi: le ipotesi richiedono stime di effetti diretti e di mediazioni testabili in modo più chiaro e preciso in uno studio sperimentale di questo tipo.

Inoltre, la decisione di svolgere gli esperimenti online è stata presa perché consente meglio la randomizzazione delle condizioni, il controllo degli stimoli, la standardizzazione della procedura e la raccolta di un campione numericamente adeguato per testare gli effetti del modello.

3.2.3 Variabili del modello e ruolo nel framework di ricerca

Sulla base delle domande di ricerca e delle ipotesi formulate, è stato sviluppato il modello concettuale di ricerca e sono stati identificati variabili indipendenti (VI), variabili dipendenti (VD) e mediatori (Med).

a) Variabili indipendenti e manipolazioni

Come variabili indipendenti (VI) manipolate, il modello assume tre costrutti che definiscono l'esperienza del consumatore rispetto a contenuti AI-generated:

1. Origine creativa del contenuto (AI vs umano), solo per lo Studio 1.
2. Disclosure dell'uso di AI (sì vs no), per entrambi gli studi.
3. Pertinenza context-aware (alta vs bassa), solo per lo Studio 2.

b) Variabili dipendenti (misure dell'efficacia dei contenuti)

Come variabili dipendenti (VD) il modello adotta due outcome, coerenti con le domande di ricerca e con la natura del DOOH:

1. Aad (atteggiamento verso l'annuncio): indicatore della valutazione dell'annuncio e dell'efficacia persuasiva del contenuto.
2. Brand memory, misurata attraverso la brand recognition (memoria assistita)

Queste due misure sono comuni ad entrambi gli studi.

c) Mediatori (meccanismi esplicativi)

Lo Studio 1 inserisce nel modello due gruppi di meccanismi esplicativi di mediazione:

1. Meccanismi percettivi di valutazione: autenticità percepita, fiducia e credibilità come driver valutativi;
2. Meccanismi difensivi: persuasion knowledge e reactance come driver difensivi.

Lo Studio 2 include meccanismi che spiegano e pertinenza: utilità e rilevanza percepita.

3.3 Partecipanti e campionamento

Popolazione di riferimento e campione

La popolazione di riferimento per entrambi gli esperimenti è composta da consumatori adulti (di età superiore ai 18 anni) con uno stile di vita che comporta una potenziale esposizione a comunicazioni DOOH in contesti urbani. L'unità di analisi è costituita dal singolo individuo come rispondente al questionario, esposto ad un singolo stimolo in una sola condizione sperimentale. Il campione effettivo utilizzato per gli esperimenti è stato reclutato online.

Criteri di inclusione ed esclusione

Per ottenere un data set di qualità (su cui poi basare l'analisi), sono stati stabiliti alcuni criteri di inclusione e di esclusione dagli esperimenti, che hanno definito i requisiti personali necessari per partecipare al sondaggio e le caratteristiche che le risposte dovevano avere per poter essere incluse nel data set finale oggetto di analisi. Entrambi i criteri di inclusione ed esclusione sono stati definiti ex ante prima dello svolgimento degli esperimenti e sono così definiti:

- a) I criteri di inclusione per la partecipazione allo studio sono: avere un'età uguale o superiore ai 18 anni, comprensione della lingua italiana e accettazione del consenso informato.
- b) I criteri di esclusione, relativi alla qualità dei dati, vedono: questionario parziale e non completo, fallimento di almeno un attention check, compilazioni eccessivamente rapide (speeders), pattern di risposta non informativi (straightlining) in combinazione con tempi anomali e risposte duplicate (quando identificabili).

Questi criteri di esclusione hanno permesso di svolgere l'analisi su un data set di qualità e di ottenere risultati reali e coerenti.

Dimensione del campione e tecnica di campionamento

La dimensione campionaria ideale è stata fissata pari a un numero di rispondenti superiore alle 200 unità, in modo da ottenere un data set di risposte sufficiente per definire i confronti pianificati tra condizioni e avere stabilità nelle stime degli effetti diretti, delle moderazioni e delle mediazioni.

Il campione finale è stato di 252 rispondenti per il primo studio e di 224 rispondenti per il secondo studio, in linea con la dimensione ideale fissata ex ante.

Il campionamento è stato di tipo non probabilistico, tipico degli esperimenti online, ma con procedure di qualità e controlli per aumentare l'affidabilità.

Procedura di reclutamento

Il reclutamento è avvenuto tramite panel online e canali universitari, in una finestra temporale dal 01 dicembre 2025 al 31 dicembre 2025. I partecipanti hanno ricevuto un link al questionario e hanno completato lo studio da remoto su dispositivo personale. Il questionario è stato configurato per permettere la compilazione sia da desktop che da mobile.

3.4 Strumenti di ricerca e raccolta dati: piattaforma, questionario e stimoli

Piattaforma di somministrazione e impostazione tecnica

Per entrambi gli studi, la raccolta dati è stata effettuata tramite un questionario (uno per ciascuno studio), implementato su Qualtrics. Per costruire i questionari, sono state sfruttate alcune impostazioni tecniche disponibili sulla piattaforma, tra cui: il randomizer per la randomizzazione delle condizioni e la memorizzazione delle condizioni tramite variabili Embedded (dati incorporati), così da semplificare la successiva analisi dei dati.

I questionari sono stati strutturati in blocchi separati sistemati in un ordine preciso e studiato, per minimizzare gli effetti di ordine e preservare la qualità delle misurazioni.

La piattaforma ha inoltre registrato il tempo totale di risposta dell'intero questionario e il completamento dello stesso, per utilizzare poi questi elementi per i controlli di qualità e il superamento dei criteri di esclusione.

Stimoli DOOH: principio di controllo e descrizione

Entrambi gli studi hanno utilizzato degli stimoli DOOH controllati. Gli stimoli (tre per lo Studio 1 e quattro per lo Studio 2) sono stati progettati come mockup di annunci DOOH, seguendo un principio di controllo: gli annunci presentati all'interno degli stimoli avevano tutti stesso brand (fittizio), stessa durata, stesso layout, stessa categoria di prodotto e stesso claim. Le uniche differenze tra i vari stimoli erano relative agli elementi collegati alle variabili manipolate: origine creativa (AI vs umano) per lo

Studio 1, pertinenza context-aware (alta vs bassa) per lo Studio 2 e disclosure (SI vs NO) per entrambi gli studi, in modo da attribuire le eventuali differenze negli outcome alle manipolazioni sperimentali.

I singoli stimoli sono stati progettati in modo da ricreare in tutto e per tutto un annuncio DOOH e le sue modalità di esposizione. Sono stati realizzati come mockup di un annuncio DOOH, in forma di clip breve (11-12 secondi) che simula una visualizzazione realistica inserita in una reale scena urbana: un maxischermo digitale collocato in un contesto urbano, con un punto di osservazione proprio di un passante. Questo formato ha permesso di simulare l'esperienza DOOH, proponendo un'esposizione breve e rapida e un contesto ambientale circostante.

Nel dettaglio, gli stimoli sono stati composti da uno sfondo urbano in cui è stato inserito uno schermo DOOH integrato nell'ambiente, al cui interno è stato posto l'annuncio pubblicitario.

La visualizzazione dello stimolo all'interno del questionario è stata impostata con riproduzione automatica e con un tempo di visione minimo prima di consentire di proseguire, così da ridurre il rischio che i rispondenti saltassero lo stimolo e per replicare al massimo le condizioni reali di esposizione DOOH.

I contenuti degli annunci proposti sono stati costruiti con headline chiara e semplice e logo/nome del brand ben visibile, in modo da fornire poche informazioni con alta leggibilità, coerentemente con le caratteristiche di un annuncio DOOH. Solo negli stimoli con condizione di disclosure "SI", è stata anche aggiunta l'etichetta "Creato con AI", posizionata in modo fisso e uguale per ogni stimolo, e pensata per essere visibile ma non dominante, così da simulare una disclosure realistica dell'uso di AI senza però trasformarla nel focus principale dell'annuncio.

Come brand sponsorizzato nell'annuncio è stato scelto un brand fittizio, in modo da evitare influenze esterne e atteggiamenti preesistenti, consentendo di attribuire le risposte esclusivamente allo stimolo e alle manipolazioni e non a eventuali bias personali verso marche note. La categoria di prodotto è stata selezionata affinché fosse una categoria neutra e a basso coinvolgimento, sempre per ridurre l'effetto di bias e preferenze personali. Tutte queste impostazioni hanno garantito che le eventuali differenze negli outcome dipendessero esclusivamente dallo stimolo e dal trattamento sperimentale.

Questionario e scale

I due questionari hanno misurato: gli outcome di efficacia (Aad e brand recognition), i mediatori del modello (autenticità, fiducia, credibilità, persuasion knowledge e reactance; utilità e rilevanza), i Manipulation checks e le variabili covariate e demografiche.

Le scale multi-item per misurare i costrutti di mediazione sono state costruite su scala Likert a 7 punti (1 = per nulla d'accordo; 7 = totalmente d'accordo); per misurare il Aad è stata utilizzata una scala semantic differential a 7 punti, mentre la brand recognition è stata misurata attraverso una domanda a risposta multipla. Le scale utilizzate sono state tutte tratte o adattate a partire da scale preesistenti consolidate. Poiché molte scale erano originariamente sviluppate in inglese, gli item sono stati tradotti e adattati in italiano con attenzione a mantenere l'equivalenza semantica. Per ridurre errori di traduzione, è stata utilizzata una procedura di traduzione e revisione di back-translation.

3.5 Procedura empirica: Studio 1

3.5.1 Disegno sperimentale, manipolazioni e randomizzazione

Lo Studio 1 ha testato le prime due domande di ricerca (RQ1 e RQ2) e il relativo set di ipotesi H1-H5. In particolare, lo studio ha testato se, in assenza di segnali espliciti, una creatività DOOH AI-generated e una creatività DOOH human-generated risultano equivalenti in termini di Aad e brand memory (RQ1; H1a e H1b) e se la presenza di disclosure “Creato con AI” riduce l'efficacia della creatività AI (RQ2; H2a e H2b) e attraverso quali meccanismi (RQ2; H3-H5).

È stato formalmente progettato con un disegno concettuale di tipo 2x2 (Origine creativa x Disclosure), con una manipolazione dell'origine creativa (AI-generated vs Human-generated) e della disclosure (SI vs NO). È stato poi implementato in modo logicamente coerente, con la randomizzazione della disclosure “Creato con AI” solo tra i partecipanti assegnati all'origine creativa AI-generated. Operativamente, ne è risultata una struttura con tre diverse condizioni:

1. C1: Human-generated x Disclosure NO
2. C2: AI-generated x Disclosure NO
3. C3: AI-generated x Disclosure SI

Relativamente a queste condizioni, è stata effettuata una randomizzazione in due step:

1. Step 1: assegnazione casuale in egual misura dell'origine creativa (AI vs Human), randomizzata per tutta la popolazione dei rispondenti.
2. Step 2: assegnazione casuale in egual misura della presenza di Disclosure (SI vs NO), randomizzata solo per quella parte di popolazione per cui l'origine creativa era AI-generated.

La piattaforma è stata impostata in modo da salvare automaticamente le variabili di condizione attraverso Embedded Data (dati incorporati), così da rendere più chiara e semplice la successiva analisi dei dati.

La struttura a più condizioni ha dato la possibilità di svolgere test di ipotesi costruiti come confronti controllati tra le varie condizioni, mantenendo le analisi statisticamente semplici.

3.5.2 Stimoli dell'esperimento e manipolazioni

La manipolazione dell'origine creativa (AI vs umana) è avvenuta attraverso la creazione di due versioni dell'annuncio (AI-generated vs human-generated) mostrato come stimolo, le quali differivano esclusivamente per il soggetto che ha generato il contenuto. Le due versioni, infatti, sono state costruite per essere il più possibile comparabili: stesso claim, stesso layout, stessa categoria e stesso brand fittizio, stesso formato video e durata. L'unica differenza tra gli stimoli riguardava esclusivamente l'origine della creatività, in modo appunto da isolare gli effetti della manipolazione. Nella versione AI, il contenuto è stato creato tramite AI generativa e rifinito manualmente per ottenere coerenza grafica; nella versione umana, il contenuto è stato creato manualmente e composto da elementi non generativi. Quest'ultima versione dell'annuncio funge da stimolo mostrato ai partecipanti assegnati alla condizione C1, mentre la versione AI è lo stimolo mostrato ai partecipanti assegnati alla condizione C2.

Prima della somministrazione definitiva degli stimoli, è stato effettuato un controllo interno per verificare che i due annunci relativi alle due creatività (AI vs umana) fossero effettivamente comparabili, evitando che eventuali differenze nella qualità estetica diventassero un elemento di confusione.

La disclosure, invece, è implementata come etichetta testuale "Creato con AI" e la sua manipolazione è avvenuta tramite la presenza o assenza dell'etichetta all'interno dell'annuncio, in base alla condizione di riferimento. L'etichetta di disclosure è coerente e identica in tutti gli stimoli in cui è presente ed ha grafica standardizzata, è sempre nella stessa posizione e ha una dimensione tale da

essere notata ma non risultare dominante. Alla condizione con “Disclosure SI” è collegata una terza versione dell’annuncio, identica a quella della C2, con l’aggiunta dell’etichetta “Creato con AI”.

Per verificare che la manipolazione della disclosure fosse effettivamente percepita come previsto, prima di inserire gli stimoli all’interno del questionario e procedere con il main study, è stato svolto un pre-test che è servito da manipulation check della disclosure. Il pre-test è stato svolto con un campione diverso rispetto a quello dello studio principale e meno numeroso (53 rispondenti nel pre-test), reclutato attraverso un panel online. Ad ognuno dei partecipanti è stato mostrato un singolo stimolo, assegnato in maniera randomica tra uno stimolo senza etichetta di disclosure e uno stimolo con etichetta di disclosure, e gli è stato poi chiesto di verificare se avessero notato la scritta, con la domanda “Hai notato la scritta “Creato con AI” all’interno dell’annuncio?”, con risposta multipla “SI” o “NO”. I risultati del pre-test hanno dimostrato che la manipolazione era percepita nel modo corretto e , quindi, gli stimoli sono stati considerati adatti per essere utilizzati nello studio principale. Il manipulation check dell’origine creativa non è stato effettuato perché ai fini di un corretto svolgimento dello studio non era necessario che i rispondenti notassero la manipolazione dell’origine creativa.

3.5.3 Struttura del questionario

La raccolta dati è stata effettuata tramite un questionario, strutturato in blocchi fissi, con lo scopo di ridurre bias di ordine e preservare la validità delle misure di memoria, misurate subito dopo l’esposizione allo stimolo. La sequenza del questionario è stata identica per tutti i rispondenti, salvo la condizione dello stimolo:

1. Consenso informato, che comprende informativa sullo scopo accademico, anonimato, durata e diritto di recesso del questionario.
2. Istruzioni generali per la compilazione del questionario, come, ad esempio, richiesta di attenzione nella risposta alle domande, visione completa dell’annuncio mostrato e non possibilità di scorrere tra le domande o tornare indietro.
3. Scenario di esposizione, per contestualizzare lo stimolo proposto.
4. Esposizione allo stimolo DOOH, con tempo minimo di visualizzazione.
5. Attention check 1, per intercettare compilazioni non attente.
6. Blocco di domande relative alla brand recognition.
7. Blocco di domande relative all’atteggiamento verso l’annuncio (Aad).

8. Domande relative ai costrutti di mediazione: autenticità, fiducia, e credibilità, poi persuasion knowledge e reactance.
9. Attention check 2.
10. Domande di controllo (covariate) e demografiche (atteggiamento verso AI, familiarità con strumenti di AI generativa, frequenza di esposizione DOOH, età, genere e titolo di studio).

3.5.4 Misure e scale utilizzate

Le misure utilizzate sono scale multi-item misurate su scala Likert a 7 punti (1 = per nulla d'accordo; 7 = totalmente d'accordo), salvo per Aad e brand recognition, e sono state adottate da scale consolidate in letteratura.

Per tutte le scale multi-item è stata valutata la coerenza e l'affidabilità interna tramite Cronbach's alpha (Cronbach, 1951).

Variabili dipendenti : Aad e brand memory

Per misurare le variabili dipendenti, sono state adottate scale classiche di misura dell'efficacia pubblicitaria.

a) L'Attitude toward the ad (Aad) è stata misurata tramite quattro semantic differential a 7 punti, adottati dal filone classico sull'Aad (MacKenzie, Lutz, & Belch, 1986; MacKenzie & Lutz, 1989).

Di seguito sono riportate le coppie bipolari utilizzate. "Questo annuncio mi sembra...":

- sgradevole / gradevole
- negativo / positivo
- brutto / bello
- non convincente / convincente

L'indice finale utilizzato per l'analisi dei risultati è rappresentato dalla media dei 4 item.

b) Per misurare la brand recognition (memoria assistita) è stata proposta una domanda a risposta multipla: "Quale brand hai visto nell'annuncio?"

Tra le opzioni: 1 corretta con il nome del brand, 3 non corrette , 1 "Non ricordo".

Studiata con la codifica: corretto = 1, altrimenti = 0.

Mediatori valutativi: autenticità, fiducia, credibilità

a) L'autenticità percepita è stata misurata con scala likert a 7 punti (1 = per nulla d'accordo; 7 = totalmente d'accordo), adottata da un filone sulla brand authenticity (Morhart et al., 2015; Napoli et al., 2014).

Qui di seguito sono riportati gli item utilizzati:

- “Questo annuncio mi sembra autentico.”
- “Questo annuncio mi sembra genuino.”
- “Questo annuncio mi sembra naturale, non artificiale.”
- “Questo annuncio mi sembra sincero.”

b) La fiducia è stata misurata con scala likert a 7 punti (1 = per nulla d'accordo; 7 = totalmente d'accordo), adottata da un filone su brand trust e credibility (Chaudhuri & Holbrook, 2001; Erdem & Swait, 2004).

Qui di seguito sono riportati gli item utilizzati:

- “Mi fido di ciò che comunica questo annuncio.”
- “Questo annuncio mi sembra affidabile.”
- “Questo annuncio mi dà una sensazione di sicurezza rispetto al brand.”

c) La credibilità percepita del messaggio è stata misurata con scala likert a 7 punti (1 = per nulla d'accordo; 7 = totalmente d'accordo), adottata da scale di message e brand credibility (Erdem & Swait, 2004).

Qui di seguito sono riportati gli item utilizzati:

- “Questo annuncio mi sembra credibile.”
- “Il messaggio di questo annuncio mi sembra plausibile.”
- “Questo annuncio mi sembra onesto.”

Mediatori difensivi: Persuasion knowledge e reactance

a) La Persuasion knowledge (PKM) è stata misurata con scala likert a 7 punti (1 = per nulla d'accordo; 7 = totalmente d'accordo), con item adottati da Friestad & Wright (1994).

Qui di seguito sono riportati gli item utilizzati:

- “Questo annuncio sta chiaramente cercando di persuadermi.”
- “Questo annuncio usa tattiche per influenzare la mia scelta.”

- “Sono consapevole delle intenzioni persuasive dietro questo annuncio.”

b) La reactance è stata misurata con scala likert a 7 punti (1 = per nulla d'accordo; 7 = totalmente d'accordo), con item adottati da Dillard & Shen (2005).

Qui di seguito sono riportati gli item utilizzati:

- “Questo annuncio mi irrita.”
- “Questo annuncio mi dà fastidio.”
- “Questo annuncio mi fa venire voglia di resistere al messaggio.”
- “Questo annuncio mi fa sentire sotto pressione nella mia libertà di scelta.”

Variabili di controllo e demografiche

Alla fine del questionario, sono state incluse delle covariate:

- Atteggiamento generale verso l'Intelligenza Artificiale, misurato attraverso 3 item con scala Likert a 7 punti.
- Familiarità con strumenti AI, misurata tramite scale Likert a 7 punti con la domanda “Quanto ti senti familiare con strumenti di AI generativa?”
- Esposizione abituale a DOOH e frequenza con cui si notano gli schermi DOOH, misurato con scala Likert a 7 punti con la domanda “Quanto spesso noti schermi pubblicitari digitali in città?”.
- Demografiche (età, genere, istruzione).

3.6 Procedura empirica: Studio 2

3.6.1 Disegno sperimentale, manipolazioni e randomizzazione

Lo studio 2 ha testato la terza domanda di ricerca (RQ3) e il relativo set di ipotesi H6-H9. In particolare, lo studio ha testato se, nel DOOH con annunci con creatività AI, un'alta pertinenza context-aware aumenta Aad e brand recognition, attraverso un aumento dell'utilità e della rilevanza percepita, e se la presenza di disclosure “Creato con AI” può diminuire questo effetto.

L'esperimento ha un disegno concettuale di tipo 2x2 (Pertinenza context-aware x Disclosure), con una manipolazione della pertinenza context-aware (alta vs bassa) e della disclosure (SI vs NO). Da

queste manipolazioni è risultata una struttura a quattro condizioni, randomizzate in parti uguali nella popolazione dei rispondenti:

1. C1: Pertinenza bassa x Disclosure NO
2. C2: Pertinenza alta x Disclosure NO
3. C3: Pertinenza bassa x Disclosure SI
4. C4: Pertinenza alta x Disclosure SI

La piattaforma è stata impostata in modo da salvare automaticamente le variabili di condizione attraverso Embedded Data (dati incorporati), così da rendere più chiara e semplice l'analisi dei dati.

La struttura a più condizioni ha dato la possibilità di svolgere test di ipotesi costruiti come confronti controllati tra le varie condizioni, mantenendo le analisi statisticamente semplici.

3.6.2 Stimoli dell'esperimento e manipolazioni

La manipolazione della pertinenza context-aware (alta vs bassa) è stata effettuata tramite la creazione di due diverse versioni dell'annuncio proposto come stimolo, una versione con contenuto ad alta pertinenza contestuale e una con contenuto a bassa pertinenza contestuale. Le due versioni sono state costruite per essere il più possibile comparabili a livello strutturale: stessa categoria e stesso brand fittizio, stesso formato video e durata; l'unica differenza riguardava esclusivamente la pertinenza del contenuto dell'annuncio rispetto al contesto, in modo così da isolarne gli effetti.

La pertinenza context-aware è stata basata su segnali contestuali non personali, in particolare meteo e periodo dell'anno.

È stato poi proposto uno scenario contestuale, uguale per tutti i partecipanti, prima dell'esposizione allo stimolo, in modo così da definire precisamente il contesto di esposizione dell'annuncio.

Nelle condizioni ad alta pertinenza, il contenuto dell'annuncio DOOH richiamava esplicitamente lo scenario precedentemente proposto; nelle condizioni a bassa pertinenza, invece, il contenuto dell'annuncio era generico e poco congruente con lo scenario proposto.

La disclosure, invece, è stata aggiunta come etichetta testuale "Creato con AI" e la sua manipolazione è avvenuta con la presenza o assenza dell'etichetta all'interno dell'annuncio, in base alla condizione di riferimento. L'etichetta di disclosure è coerente e identica in tutti gli stimoli in cui è presente ed ha grafica standardizzata, è sempre nella stessa posizione e ha una dimensione tale da essere notata ma non dominante.

Da queste manipolazioni sono risultati quattro stimoli diversi; ognuno di questi stimoli era associato a una delle quattro condizioni e quindi mostrato esclusivamente ai partecipanti associati a quella condizione.

Per verificare che le manipolazioni della pertinenza context-aware e della disclosure fossero effettivamente percepite come previsto, prima di inserire gli stimoli all'interno del questionario e procedere con il main study, è stato svolto due pre-test serviti da manipulation check della disclosure e della pertinenza. I pre-test sono stati svolti con un campione diverso rispetto a quello dello studio principale e meno numeroso (47 rispondenti nel pre-test per il check della disclosure, 51 in quello per la pertinenza), reclutato attraverso un panel online. Nel primo pre-test, ad ognuno dei partecipanti è stato mostrato un singolo stimolo, assegnato in maniera randomica tra uno stimolo senza etichetta di disclosure e uno stimolo con etichetta di disclosure, e gli è stato poi chiesto di verificare se avessero notato la scritta, con la domanda "Hai notato la scritta "Creato con AI" all'interno dell'annuncio?", con risposta multipla "SI" o "NO". Il secondo pre-test ha avuto stessa struttura; ad ogni rispondente è stato mostrato un solo stimolo assegnato in maniera randomica tra uno con basso livello di pertinenza e uno con alto livello di pertinenza. I risultati dei pre-test hanno dimostrato che entrambe le manipolazioni erano percepite nel modo corretto e, quindi, gli stimoli sono stati considerati adatti per essere utilizzati nello studio principale.

3.6.3 Struttura del questionario

La struttura del questionario è stata analoga a quella dello Studio 1.

La raccolta dati è stata effettuata tramite un questionario, strutturato in blocchi fissi, con lo scopo di ridurre bias di ordine e preservare la validità delle misure di memoria, misurate subito dopo l'esposizione allo stimolo. La sequenza del questionario è stata identica per tutti i rispondenti, salvo la condizione dello stimolo:

1. Consenso informato, che comprende informativa sullo scopo accademico, anonimato, durata e diritto di recesso del questionario.
2. Istruzioni generali per la compilazione del questionario, come, ad esempio, richiesta di attenzione nella risposta alle domande, visione completa dell'annuncio mostrato e impossibilità di scorrere tra le domande.
3. Scenario contestuale specifico e dettagliato, per definire precisamente il contesto di esposizione.
4. Esposizione allo stimolo DOOH, con tempo minimo di visualizzazione.

5. Attention check 1, per intercettare compilazioni non attente.
6. Blocco di domande relativo alla brand recognition.
7. Blocco di domande relativo all'atteggiamento verso l'annuncio (Aad).
8. Domande relative ai costrutti di mediazione: utilità e rilevanza.
9. Attention check 2.
10. Domande di controllo (covariate) e demografiche (atteggiamento verso AI, familiarità con strumenti di AI generativa, frequenza di esposizione DOOH, età, genere e titolo di studio).

3.6.4 Misure e scale utilizzate

Le misure utilizzate sono scale multi-item misurate su scala Likert a 7 punti (1 = per nulla d'accordo; 7 = totalmente d'accordo), salvo per Aad e brand recognition, e sono state adottate da scale consolidate in letteratura.

Per tutte le scale multi-item è stata valutata la coerenza e l'affidabilità interna tramite Cronbach's alpha (Cronbach, 1951).

Le Variabili dipendenti (Aad, recognition, recall) sono state misurate in modo identico allo Studio 1 per avere comparabilità tra gli studi.

Variabili dipendenti: Aad e brand memory

Per misurare le variabili dipendenti, sono state adottate scale classiche di misura dell'efficacia pubblicitaria.

a) L'Attitude toward the ad (Aad) è stata misurata tramite quattro semantic differential a 7 punti, adattati da un filone classico sull'Aad (MacKenzie, Lutz, & Belch, 1986; MacKenzie & Lutz, 1989). Di seguito sono riportate le coppie bipolari utilizzate. "Questo annuncio mi sembra...":

- sgradevole / gradevole
- negativo / positivo
- brutto / bello
- non convincente / convincente

Il valore finale utilizzato per l'analisi dei risultati è rappresentato dalla media dei 4 item.

b) Per misurare la brand recognition (memoria assistita) è stata proposta una domanda a risposta multipla: “Quale brand hai visto nell’annuncio?”

Tra le opzioni: 1 corretta con il nome del brand, 3 non corrette , 1 “Non ricordo”.

Studiata con la codifica: corretto = 1, altrimenti = 0.

Mediatori di valore: utilità e rilevanza

L'utilità e la rilevanza percepita sono state misurate insieme con una scala likert a 7 punti (1 = per nulla d'accordo; 7 = totalmente d'accordo), adattata da Ducoffe (1996) e Davis (1989).

Qui di seguito sono riportati gli item utilizzati:

- “Questo annuncio è rilevante per la situazione descritta.”
- “Questo annuncio è utile nel contesto in cui appare.”
- “Questo annuncio mi sembra pertinente rispetto al momento/luogo.”
- “Questo annuncio ha senso qui e ora.”

Variabili di controllo e demografiche

Alla fine del questionario, sono state incluse delle covariate:

- Atteggiamento generale verso l'Intelligenza Artificiale, misurato attraverso 3 item con scala Likert a 7 punti.
- Familiarità con strumenti AI, misurata tramite scale Likert a 7 punti con la domanda “Quanto ti senti familiare con strumenti di AI generativa?”
- Esposizione abituale a DOOH e frequenza con cui si notano gli schermi DOOH, misurato con scala Likert a 7 punti con la domanda “Quanto spesso noti schermi pubblicitari digitali in città?”.
- Demografiche (età, genere, istruzione).

3.7 Analisi dei dati

L'analisi è stata condotta su SPSS, utilizzato per tutto il processo di analisi.

I dati da analizzare sono stati esportati dalla piattaforma Qualtrics e importati poi su SPSS; all'interno dei dati la piattaforma ha salvato le risposte ai questionari, le variabili utili per i controlli di qualità

(completamento sì/no e tempi di compilazione) e le variabili di condizione, ovvero gli Embedded Data, fondamentali per collegare i rispondenti alla loro condizione sperimentale.

3.7.1 Struttura e preparazione del dataset

Struttura dei dataset

Sono state svolte due analisi separate, una per ogni studio, utilizzando due dataset distinti, uno per lo Studio 1 e uno per lo Studio 2.

I dataset sono stati mantenuti nella struttura wide, con ogni riga che corrispondeva ad un singolo partecipante (rispondente) e con varie colonne relative a:

- Variabili sperimentali/variabili indipendenti (condizioni);
- Outcome/variabili dipendenti (Aad, brand recognition);
- Mediatori;
- Covariate e demografiche.

L'unità di analisi è stata, quindi, il singolo individuo rispondente al questionario.

Pulizia del dataset e criteri di qualità

Come prima cosa, sono stati puliti i dataset applicando i criteri di esclusioni definiti ex ante, per ottenere un dataset di qualità su cui svolgere l'analisi. In particolare, sono stati esclusi:

- I questionari incompleti (risposte parziali o non terminate);
- Le risposte dei partecipanti che hanno fallito almeno un attention check, perché indicatore di non attenzione);
- Le compilazioni troppo rapide, guardando i tempi di completamento;
- Le risposte duplicate, quando identificabili.

Dopo l'applicazione dei criteri di esclusione, il campione finale dello Studio 1 è stato di 254 unità (rispondenti), mentre il campione finale dello Studio 2 è stato di 236 unità (rispondenti).

3.7.2 Codifica delle variabili e costruzione degli indici

Variabili di condizione (variabili indipendenti/manipolate)

Per entrambi gli studi e in entrambi i dataset, le condizioni sperimentali sono state recuperate dalle variabili di Embedded Data registrate automaticamente da Qualtrics.

Per lo studio 1, nel dataset è stata creata una variabile di condizione (COND_S1) a 3 livelli, una per condizione. Per lo Studio 2, nel dataset sono state create due variabili fattoriali binarie PERTINENZA (0 = bassa; 1 = alta) e DISCLOSURE (0 = NO; 1 = SI) e una variabile di condizione (COND_S2) a quattro livelli.

Variabili dipendenti (outcome)

Gli outcome sono uguali in entrambi gli studi. L'indice finale per studiare l'Aad è stato costruito come media aritmetica dei 4 item della scala, mentre la brand recognition è stata codificata come variabile binaria (1 se risposta corretto; 0 se altrimenti).

Mediatori e indici composti degli item

Per ogni mediatore, è stato costruito un indice finale usato per studiare il costrutto, creato come media degli item della relativa scala.

3.7.3 Analisi preliminari

Prima dei test d'ipotesi, sono state effettuate delle analisi preliminari.

Statistiche descrittive

Come primo passo sono state prodotte le statistiche descrittive per tutte le variabili principali, in modo da descrivere il campione ed individuare eventuali anomalie:

- Media, deviazione standard e range, per le variabili continue;
- Frequenze e percentuali, per le variabili binarie.

Affidabilità interna delle scale

Per tutte le scale multi-item è stata condotta un'analisi di affidabilità, valutando la coerenza interna tramite Cronbach's alpha, stabilendo una buona coerenza interna per valori di α compresi tra 0.8 e 0.9 e un'eccellente coerenza interna per valori di α superiori o uguali a 0.9.

3.7.4 Strategia di analisi inferenziale e test delle ipotesi

Tutte le analisi inferenziali sono state condotte adottando un livello di significatività pari a $\alpha = 0.05$. La scelta delle tecniche di analisi si è basata sulla natura dell'outcome:

- a) Aad e mediatori (misurati con scale a 7 punti) sono stati trattati come variabili continue e, quindi, analizzati attraverso t-test e regressione lineare;
- b) La brand recognition è una variabile binaria (0/1) e, quindi, è stata analizzata attraverso test chi-quadrato e regressione logistica.

3.7.5 Analisi Studio 1: test delle ipotesi H1–H5

Allo studio sono collegate tre condizioni sperimentali:

- C1: Human-generated x Disclosure NO
- C2: Human-generated x Disclosure SI
- C3: AI-generated x Disclosure SI

Il set di ipotesi relativo è stato testato tramite confronti pianificati tra coppie di condizioni: C1 vs C2 per H1a e H1b; C2 vs C3 per H2-H5.

L'analisi è stata condotta sul campione finale dello Studio 1, pari a 252 rispondenti.

Test di H1a e H1b

H1a ipotizza che, in assenza di segnali espliciti, un contenuto con origine creativa AI-generated e un contenuto con origine creativa human-generated non differiscono in termini di Aad.

Per testare H1a è stato condotto un confronto pianificato tra la condizione C1 (Human-generated x Disclosure NO) e la condizione C2 (AI-generated x Disclosure NO) su Aad.

Poiché Aad è una variabile continua, per testare l'ipotesi è stato utilizzato un t-test per campioni indipendenti (con Aad come variabile dipendente e la condizione come variabile indipendente di raggruppamento).

H1b ipotizza che, in assenza di segnali espliciti, un contenuto con origine creativa AI-generated e un contenuto con origine creativa human-generated non differiscono in termini di brand recognition.

Per testare H1b è stato condotto un confronto tra la condizione C1 e la condizione C2 su brand recognition.

Poiché la brand recognition è una variabile binaria, H1b è stata testata attraverso una Crosstabs e un test chi-quadrato.

Test di H2a e H2b

H2a e H2b ipotizzano che, in presenza di disclosure “Creato con AI”, il contenuto di origine creativa AI-generated è meno efficace su Aad e brand recognition rispetto alla condizione con assenza di disclosure. H2a e H2b sono state testate con confronto pianificato tra la condizione C2 (AI-generated × Disclosure NO) e la condizione C3 (AI-generated × Disclosure SI).

Per il confronto tra le due condizioni rispetto a Aad (test di H2a) è stato utilizzato un t-test a campioni indipendenti, con Aad come outcome e la condizione come variabile di raggruppamento.

Per il confronto tra le due condizioni rispetto a brand recognition (test di H2b) è stata utilizzata una crosstabs e un test chi-quadrato.

Test di H3

H3 ipotizza che la presenza di disclosure riduca autenticità percepita, fiducia e credibilità rispetto alla condizione senza disclosure.

Per testare H3 è stato effettuato un confronto tra la condizione C2 e la condizione C3 su autenticità, fiducia e credibilità.

Dal momento che si tratta di indici continui, l'analisi è stata implementata come t-test a campioni indipendenti su ciascuno dei tre mediatori, attraverso tre confronti separati.

Test di H4

L'ipotesi H4, come le altre ipotesi di mediazione, è stata testata tramite un modello di regressione, implementato in SPSS mediante la macro PROCESS, utile perché consente di stimare:

- a) L'effetto della variabile manipolata sul mediatore (path a)
- b) L'effetto del mediatore sull'outcome, controllando la variabile controllata (path b)
- c) L'effetto diretto della variabile manipolata sull'outcome (path c') e l'effetto totale (path c)
- d) L'effetto diretto (axb) con stima tramite bootstrap.

Tutte le mediazioni sono state testate con il modello di mediazione PROCESS Model 4. L'intervallo di confidenza per il Bootstrap è stato definito al 95% e l'effetto è stato considerato supportato se l'intervallo di confidenza bootstrap non includeva lo zero.

H4 ipotizza che l'effetto negativo della presenza di disclosure su Aad e brand memory è mediato da autenticità, fiducia e credibilità.

Per testare l'ipotesi, sono stati stimati tre modelli di mediazione semplice, uno per ogni mediatore, in cui la variabile indipendente (X) influenza il mediatore (M), che a sua volta influenza l'outcome (Y).

Il modello di mediazione in è:

- X = Disclosure (codificata come: 0 in C2; 1 in C3)
- M(1) = autenticità; M(2) = fiducia; M(3) = credibilità
- Y = Aad

L'equazione finale è stata una regressione lineare.

Test di H5

H5 ipotizza che la presenza di disclosure aumenti la persuasion knowledge e la reactance rispetto alla condizione senza disclosure.

Per testare H5 è stato effettuato un confronto tra la condizione C2 e la condizione C3 su persuasion knowledge e reactance.

Dal momento che si tratta di indici continui, l'analisi è stata implementata come t-test a campioni indipendenti su entrambi i mediatori, attraverso due confronti separati.

3.7.6 Analisi Studio 2: ipotesi H6–H9

Allo studio sono collegate quattro condizioni sperimentali:

- C1 (pertinenza bassa × disclosure NO)
- C2 (pertinenza alta × disclosure NO)
- C3 (pertinenza bassa × disclosure SÌ)
- C4 (pertinenza alta × disclosure SÌ).

Per testare le relative ipotesi H6–H9, sono stati effettuati dei confronti pianificati tra coppie di condizioni.

Test di H6a e H6b

H6a e H6b ipotizzano che un annuncio con alta pertinenza context-aware genera maggiore Aad e brand recognition rispetto ad un annuncio con bassa pertinenza.

L'ipotesi H6a è stata testata tramite un confronto pianificato tra la condizione C1 (bassa pertinenza x Disclosure NO) e la condizione C2 (alta pertinenza x Disclosure NO) su Aad.

Poiché l'outcome è una variabile continua, il confronto è stato testato con un t-test per campioni indipendenti, con Aad come test variable e le due condizioni come grouping variable.

L'ipotesi H6b è stata testata tramite lo stesso confronto pianificato della H6a, ma con brand recognition come outcome.

Poiché la brand memory è una variabile binaria (codificata 0/1), il confronto è stato stimato come confronto tra proporzioni, con Crosstabs e test Chi-quadrato.

Test di H7

La H7 ipotizza che un annuncio con alta pertinenza context-aware aumenta utilità e rilevanza percepita.

Come per le ipotesi precedenti, anche la H7 è stata testata attraverso un confronto pianificato tra le due condizioni C1 e C2 su utilità e rilevanza.

Poiché l'indice utilità/rilevanza è continuo, il test utilizzato per testare i confronti è stato un t-test per campioni indipendenti, con utilità/rilevanza come outcome e le due condizioni come variabile di raggruppamento.

Test di H8a e H8b

H8a e H8b ipotizzano che l'effetto dell'alta pertinenza context-aware su Aad e brand memory sia mediato positivamente da utilità e rilevanza percepita.

Per testare H8a è stato stimato un modello di mediazione semplice, in cui:

- X = pertinenza (codificata come: 0 = bassa; 1 = alta)
- M = utilità/rilevanza
- Y = Aad

Nel modello di mediazione semplice stimato: la pertinenza predice utilità/rilevanza (path a); utilità/rilevanza predice l'outcome, controllando la pertinenza (path b); si stima l'effetto indiretto (a x b) tramite bootstrap con intervalli di confidenza CI 95%.

Allo stesso modo, per testare H8b è stato stimato un modello di mediazione semplice, in cui:

- X = pertinenza (codificata come: 0 = bassa; 1 = alta)
- M = utilità/rilevanza
- Y = brand recognition

Nel modello di mediazione semplice stimato: la pertinenza predice utilità/rilevanza (path a); utilità/rilevanza predice l'outcome, controllando la pertinenza (path b); si stima l'effetto indiretto (a x b) tramite bootstrap con intervalli di confidenza CI 95%.

Test di H9a e H9b

H9a e H9b ipotizzano che la presenza di disclosure attenua l'effetto positivo dell'alta pertinenza context-aware su Aad e brand memory.

Per testare H9a e H9b sono state interpretate le simple effects (effetti semplici), ovvero gli effetti dell'alta pertinenza quando non c'è disclosure (disclosure = 0) e, separatamente, gli effetti dell'alta pertinenza quando è presente disclosure (disclosure = 1).

Per farlo, è stato implementato un confronto pianificato tra le due condizioni ad alta pertinenza: C2 e C4.

Per l'ipotesi H9a, il confronto pianificato su Aad è stato effettuato come t-test per campioni indipendenti, dal momento che Aad è una variabile continua (Aad come test variable; COND-S2 come variabile di raggruppamento).

Per H9b, il confronto pianificato su brand memory è stato effettuato utilizzando Crosstabs e test Chi-quadrato, dal momento che si tratta di una variabile binaria.

Capitolo IV. Risultati

4.1 Risultati analisi preliminari – Studio 1

Composizione del campione - Studio 1

Il campione finale analizzato nello Studio 1 è stato di 252 rispondenti, composto per il 47,6% da donne (120 unità) e per il restante 52,4% da uomini (132 unità), con un'età media del campione di 32 anni. 128 rispondenti sono stati assegnati alla condizione C1, 64 alla C2 e 64 alla C3.

Analisi di affidabilità Cronbach's alpha – Studio 1

Per l'analisi di affidabilità, è stata considerata una buona coerenza interna per valori di alpha compresi tra 0.8 e 0.9 e un'eccellente coerenza interna per valori uguali o superiori a 0.9.

Tutte le scale hanno mostrato valori di alpha di Cronbach superiori a 0.9 (Aad: $\alpha = 0.911$; Autenticità: $\alpha = 0.972$; Fiducia: $\alpha = 0.926$; Credibilità: $\alpha = 0.965$; PKM: $\alpha = 0.907$; Reactance: $\alpha = 0.975$), dimostrando eccellente affidabilità e coerenza interna; per questo, sono stati mantenuti tutti gli item originali in tutte le scale.

4.2 Risultati test d'ipotesi – Studio 1

H1a e H1b

Per testare H1a è stato effettuato un t-test per campioni indipendenti confrontando la condizione C1 e la condizione C2 sull'Aad. Il confronto non risulta statisticamente significativo ($t(121) = -1,690$; $p = 0,094$), evidenziando che non c'è una differenza tra i rispondenti che hanno osservato un annuncio Human-generated ($M = 5,77$; $SD = 0,65$) e i rispondenti che hanno osservato un'annuncio AI-generated ($M = 5,95$; $SD = 0,57$). In base a questi risultati, H1a, che ipotizzava assenza di differenze, è confermata.

Per testare H1b è stata condotta una crosstabs con test chi-quadrato confrontando C1 e C2 sulla brand recognition. Il test ha evidenziato un'associazione statisticamente non significativa ($\chi^2(1) = 3,41$, $p = 0,065$), mostrando che non c'è una differenza tra i rispondenti che hanno osservato un annuncio Human-generated (91% di risposte corrette) e rispondenti che hanno osservato un'annuncio AI-generated (87,5% di risposte corrette). In base a questi risultati, H1b è confermata.

H2a e H2b

Per testare H2a è stato effettuato un t-test per campioni indipendenti per confrontare le condizioni C2 e C3 sull'Aad. I risultati hanno mostrato che il confronto è statisticamente significativo ($t(122) = 20,03$; $p < 0,001$) e che c'è quindi una differenza significativa tra i rispondenti esposti a un annuncio senza disclosure ($M = 5,95$; $SD = 0,57$) e i rispondenti esposti a un annuncio con disclosure ($M = 2,75$; $SD = 1,13$) sull'Aad. In base a questi risultati, H2a è confermata.

Per testare H2b è stata condotta una crosstabs con test chi-quadrato confrontando C2 e C3 sulla brand recognition. Il test evidenzia un'associazione statisticamente significativa ($\chi^2(1) = 17,55$; $p < 0,001$), evidenziando una differenza tra i rispondenti che hanno osservato un annuncio Human-generated (87,5% di risposte corrette) e rispondenti che hanno osservato un'annuncio AI-generated (53,3% di risposte corrette) sulla brand recognition. In base a questi risultati, H1b è confermata.

H3

Per testare H3 sono stati condotti tre t-test per campioni indipendenti confrontando C2 e C3 su, rispettivamente, autenticità, fiducia e credibilità. Quanto all'autenticità, i risultati hanno mostrato una differenza statisticamente significativa ($t(122) = 42,7$; $p < 0,001$) tra i rispondenti esposti a un annuncio senza disclosure ($M = 5,44$; $SD = 0,42$) e i rispondenti esposti a un annuncio con disclosure ($M = 1,72$; $SD = 0,54$). Anche per la fiducia è stato riscontrato un confronto statisticamente significativo ($t(122) = 41,1$; $p < 0,001$) indicando una differenza tra i rispondenti esposti a un annuncio senza disclosure ($M = 5,75$; $SD = 0,54$) e i rispondenti esposti a un annuncio con disclosure ($M = 1,84$; $SD = 0,52$). Infine, anche il test sulla credibilità è risultato statisticamente significativo ($t(122) = 41,1$; $p < 0,001$) indicando anche qui una differenza tra i rispondenti esposti a un annuncio senza disclosure ($M = 6,04$; $SD = 0,62$) e i rispondenti esposti a un annuncio con disclosure ($M = 2,09$; $SD = 0,61$). Sulla base di questi tre risultati, H3 è confermata.

H4

Per testare H4 sono state stimate tre mediazioni semplici con X = disclosure, M = autenticità, fiducia, pertinenza (rispettivamente), Y = Aad. Nel modello con l'autenticità, l'effetto di X su M è risultato negativo e significativo (coeff = -3,72; SE = 0,87; $p < 0,001$), mentre l'effetto di M su Y è risultato positivo e significativo (coeff = 0,73; SE = 0,15; $p < 0,001$). L'effetto indiretto (Effect = -2,73, con IC bootstrap 95% (-3,54; -2,04) che non contiene lo zero) è significativo e conferma la mediazione. Nel modello con la fiducia, l'effetto di X su M è risultato negativo e significativo (coeff = -3,91; SE = 0,09; $p < 0,001$), mentre l'effetto di M su Y è risultato positivo e significativo (coeff = 0,97; SE = 0,13; $p < 0,001$). L'effetto indiretto (Effect = -3,77, con IC bootstrap 95% (-4,78; -2,84) che non contiene lo zero) è significativo e conferma la mediazione.

Nel modello con la credibilità, l'effetto di X su M è risultato negativo e significativo (coeff = -3,95; SE = 0,11; $p < 0,001$), mentre l'effetto di M su Y è risultato positivo e significativo (coeff = 0,26; SE = 0,13; $p = 0,049$). L'effetto indiretto (Effect = -1,02, con IC bootstrap 95% (-1,79; -0,03) che non contiene lo zero) è significativo e conferma la mediazione. Questi risultati insieme confermano l'ipotesi H4

H5

Per testare H5 sono stati condotti due t-test per campioni indipendenti confrontando C2 e C3 su, rispettivamente, persuasion knowledge e reactance. Per gli effetti sulla persuasion knowledge, i risultati hanno mostrato una differenza statisticamente significativa ($t(122) = -19,129$; $p < 0,001$) tra i rispondenti esposti a un annuncio senza disclosure (M = 3,18; SD = 0,82) e i rispondenti esposti a un annuncio con disclosure (M = 5,98; SD = 0,79). Anche per la reactance è stato riscontrato un confronto statisticamente significativo ($t(122) = -34,44$; $p < 0,001$) indicando una differenza tra i rispondenti esposti a un annuncio senza disclosure (M = 1,28; SD = 0,41) e i rispondenti esposti a un annuncio con disclosure (M = 5,55; SD = 0,89). Sulla base di questi due risultati, possiamo dire che H5 è confermata.

4.3 Risultati analisi preliminari – Studio 2

Composizione del campione - Studio 2

Il campione finale analizzato nello Studio 2 è stato di 224 rispondenti, composto per il 53,6% da donne (120 unità) e per il restante 46,4% da uomini (104 unità), con un'età media del campione di 34 anni. 56 rispondenti sono stati assegnati alla condizione C1, 48 alla C2 e 64 alla C3, 56 alla C4

Analisi di affidabilità Cronbach's alpha – Studio 2

Le scale hanno mostrato valori di alpha di Cronbach superiori a 0.9 (Aad: $\alpha = 0.960$; Utilità/rilevanza: $\alpha = 0.997$), dimostrando eccellente affidabilità e coerenza interna; per questo, sono stati mantenuti tutti gli item originali in entrambe le scale.

4.4 Risultati test d'ipotesi – Studio 2

H6a e H6b

Per testare H6a, è stato fatto un confronto tra le due condizioni C1 e C2 sull'Aad attraverso un t-test per campioni indipendenti. I risultati mostrano che il confronto è statisticamente significativo ($t(102) = -41,3$; $p < 0,001$), evidenziando una differenza tra i rispondenti che hanno osservato un annuncio a bassa pertinenza context-aware ($M = 2,92$; $SD = 0,48$) e i rispondenti che hanno osservato un annuncio a alta pertinenza context-aware ($M = 6,5$; $SD = 0,39$). In base a questi risultati, H6a è confermata.

Per testare H6b, è stata condotta una crosstabs con test chi-quadrato confrontando C1 e C2 sulla brand recognition. Il test evidenzia un'associazione statisticamente significativa ($\chi^2(1) = 26,74$; $p < 0,001$), evidenziando una differenza tra i rispondenti che hanno osservato un annuncio con bassa pertinenza (57,1% di risposte corrette) e rispondenti che hanno osservato un'annuncio con alta pertinenza (100% di risposte corrette). In base a questi risultati, H6b è confermata.

H7

Per testare H7 è stato condotto un t-test per campioni indipendenti confrontando le condizioni C1 e C2 sull'indice utilità/rilevanza. I risultati hanno mostrato un confronto statisticamente significativo ($t(102) = -70,43$; $p < 0,001$) e una differenza tra i rispondenti esposti a un annuncio con bassa pertinenza ($M = 1,5$; $SD = 0,45$) e i rispondenti esposti a un annuncio con alta pertinenza ($M = 6,70$; $SD = 0,27$). In base ai risultati, H7 è confermata.

H8a e H8b

Per testare H8a è stata stimata una mediazione semplice con X = pertinenza, M = utilità/rilevanza, Y = Aad. L'effetto di X su M è risultato significativo (coeff = 5,26; $SE = 0,52$; $p < 0,001$), come anche l'effetto di M su Y (coeff = 1,08; $SE = 0,22$; $p < 0,001$). L'effetto indiretto (Effect = 5,71, con IC bootstrap 95% (3,59; 7,64) che non contiene lo zero) è significativo e conferma la mediazione. In base a questo, H8a è confermata.

Per testare H8b è stata stimata una mediazione semplice con X = pertinenza, M = utilità/rilevanza, Y = brand recognition. L'effetto di X su M è risultato significativo (coeff = 5,26; $SE = 0,52$; $p < 0,001$), come anche l'effetto di M su Y (coeff = 3,44; $SE = 0,51$; $p < 0,001$). L'effetto indiretto (Effect = 18,11, con IC bootstrap 95% (12,63; 23,49) che non contiene lo zero) è significativo e conferma la mediazione. In base a questo, H8b è confermata.

H9a e H9b

Per testare H9a sono state confrontate le condizioni C2 e C4 su Aad attraverso un t-test per campioni indipendenti. Il confronto è risultato statisticamente significativo ($t(102) = 53,71$; $p < 0,001$) e ha mostrato una differenza significativa tra i rispondenti esposti a un annuncio senza disclosure ($M = 6,50$; $SD = 0,39$) e i rispondenti esposti a un annuncio con disclosure ($M = 2,96$; $SD = 0,28$). In base ai risultati, possiamo affermare che H9a è confermata.

Per testare H9b, è stata condotta una crosstabs con test chi-quadrato confrontando C2 e C4 sulla brand recognition. Il test evidenzia un'associazione statisticamente significativa ($\chi^2(1) = 26,74$; $p < 0,001$), evidenziando una differenza tra i rispondenti che hanno osservato un annuncio senza disclosure (100% di risposte corrette) e rispondenti che hanno osservato un annuncio con disclosure (57,1% di risposte corrette). In base a questi risultati, anche H9b è confermata.

Capitolo V. Discussione e Conclusioni

L'obiettivo della presente tesi è stato quello di studiare se e come l'utilizzo di Generative AI nella creazione di contenuti DOOH potesse influenzare l'efficacia pubblicitaria del mezzo. In particolare, sono stati testati gli effetti di origine creativa, disclosure e pertinenza context-aware su due outcome di efficacia, l'Aad e la brand recognition, e i meccanismi sottostanti.

Per rispondere alle domande di ricerca e testare le ipotesi del modello, è stato adottato un approccio di ricerca quantitativo e sperimentale, articolato in due studi a cui hanno fatto riferimento due esperimenti separati. Uno studio ha trattato gli effetti di origine creativa e disclosure, l'altro gli effetti di pertinenza context-aware e disclosure.

In questo capitolo, verranno discussi e interpretati i risultati ottenuti attraverso l'analisi dei dati; inoltre, verranno presentati i contributi teorici e le implicazioni manageriali supportati dai risultati della tesi, le limitazioni del presente lavoro e alcune possibili direzioni di ricerca futura riguardo al tema trattato.

5.1 Discussione e interpretazione dei risultati

Tutte le ipotesi del modello sono state accettate attraverso i risultati dell'analisi dei dati, confermando tutto ciò che era stato precedentemente ipotizzato.

Nello Studio 1, il test di H1 ha confermato che, in assenza di elementi che non esplicitano direttamente l'origine creativa, l'efficacia pubblicitaria di un contenuto DOOH AI-generated e di un contenuto DOOH human-generated non è significativamente diversa in termini di Aad e brand recognition. Questo risultato fa capire che, in un contesto in cui l'attenzione e l'esposizione al messaggio sono limitate e le percezioni dei consumatori si basano sulla prima impressione, la creatività umana e la creatività AI, se non esplicitate, sono difficili da distinguere e, quindi, portano a effetti di efficacia simili.

Un cambiamento di questi effetti si ha quando l'origine creativa AI-generated è esplicitata attraverso un'etichetta di disclosure "Creato con AI", come mostrato dallo studio delle altre ipotesi di ricerca. Il test di H2 ha dimostrato che, quando subentra la disclosure dell'uso di AI, gli effetti sull'efficacia pubblicitaria cambiano, mostrando che un annuncio con etichetta "Creato con AI" ha degli effetti di efficacia in termini di Aad e brand recognition peggiori rispetto allo stesso identico annuncio senza

etichetta. La disclosure, rendendo esplicito e immediatamente visibile l'utilizzo di AI, sposta l'attenzione dei consumatori su altri elementi, cambiandone la valutazione in negativo. Questo cambiamento è spiegato attraverso alcuni meccanismi, testati con le ipotesi H3-H5, che spiegano il perché la presenza di disclosure riduca l'efficacia del DOOH. I test di queste ipotesi hanno innanzitutto confermato che la presenza di disclosure fa diminuire le valutazioni dei consumatori riguardo a autenticità percepita, fiducia e credibilità del contenuto e che proprio questi meccanismi mediano l'effetto negativo della disclosure sull'efficacia in termini di Aad; inoltre, i risultati hanno anche mostrato che la presenza di disclosure fa aumentare la consapevolezza dei consumatori circa gli intenti persuasivi dell'annuncio e incrementa le loro azioni difensive e di resistenza nei confronti dello stesso, suggerendo potenzialmente un altro motivo per cui la disclosure diminuisca l'efficacia del contenuto.

Lo Studio 2 ha analizzato un altro elemento fondamentale legato all'utilizzo di AI, ovvero la sua capacità di creare contenuti contestualmente pertinenti. In particolare, ha studiato se contenuti ad alta e a bassa pertinenza abbiano degli effetti diversi in termini di efficacia e attraverso quali meccanismi. I risultati del test di H6 hanno mostrato che un annuncio ad alta pertinenza context-aware è più efficace in termini di Aad e brand recognition rispetto ad un annuncio con bassa pertinenza context-aware; inoltre, i test di H7 e H8 hanno fatto emergere che l'alta pertinenza in un annuncio aumenta l'utilità del contenuto e la rilevanza percepita e che proprio questi due fattori spiegano, attraverso mediazione, gli effetti positivi dell'alta pertinenza context-aware su Aad e brand recognition, indicando che, se un messaggio è ritenuto utile e rilevante, è più facile per i consumatori ricordarlo e valutarlo in maniera positiva. Infine, è stato analizzato il ruolo della disclosure in questo contesto, confermando che un contenuto context-aware in cui è presente l'etichetta "Creato con AI" ha degli effetti di efficacia in termini di Aad e brand recognition peggiori rispetto allo stesso contenuto senza etichetta. Questo fa capire che, quando l'utilizzo di AI è esplicitato tramite disclosure, l'effetto positivo di un'alta pertinenza è indebolito.

5.2 Contributi teorici della ricerca

I risultati della presente tesi contribuiscono sia alla ricerca sull'AI nella comunicazione di marketing, sia alla ricerca sul DOOH, ampliandole attraverso vari contributi.

La maggior parte degli studi sull'utilizzo di AI nella comunicazione e nell'advertising è concentrata in contesti digitali; la presente ricerca integra questo filone studiando le applicazioni dell'AI nel

DOOH, un contesto ambientale e fisico. Allo stesso modo, contribuisce anche alla ricerca sul DOOH, da sempre frammentata e sottosviluppata, integrandola con delle evidenze riguardo la sua efficacia in relazione all'utilizzo di AI per la creazione di contenuti.

Un altro contributo riguarda la letteratura in materia di disclosure e, in particolare, la sua natura di variabile ambivalente. Una parte della letteratura tratta la disclosure come un segnale di trasparenza e correttezza, un'altra parte, invece, la considera un elemento che amplifica la percezione da parte dei consumatori dell'intento persuasivo di un messaggio e che, quindi, peggiora la loro valutazione. La tesi contribuisce a questa discussione mostrando che, nel DOOH e per la forma specifica di disclosure testata ("Creato con AI"), l'effetto prevalente della presenza di disclosure è un effetto negativo che può diminuire l'efficacia pubblicitaria e attivare reazioni difensive e resistenza.

Infine, la tesi contribuisce anche alla letteratura in materia di pertinenza context-aware, mostrando che, quando la pertinenza è basata su segnali contestuali non personali e viene percepita come utile e rilevante, può essere un elemento che migliora l'efficacia pubblicitaria della comunicazione.

5.3 Implicazioni manageriali della ricerca

I risultati ottenuti nella presente tesi offrono alcuni spunti e indicazioni manageriali utili per lo sviluppo di strategie di comunicazione e campagne di advertising che vogliono implementare anche il DOOH nel proprio media mix, in particolare riguardo a quando usare l'AI nel DOOH, come usarla, come gestire la disclosure e come progettare la pertinenza context-aware.

Una prima implicazione manageriale è che, quando non c'è l'obbligo o la volontà di inserire un'etichetta di disclosure, la Generative AI può essere un buono strumento di produzione creativa nel DOOH, dal momento che la sua capacità di produrre contenuti è comparabile con quella umana, sia in termini visivi che in termini di efficacia. L'utilizzo di AI permetterebbe sia di diminuire costi e tempi, che di avere più versioni e adattamenti creativi.

Quando, invece, è necessario inserire la disclosure dell'utilizzo di AI all'interno del contenuto, la scelta sull'utilizzare o meno l'AI deve essere più ponderata, dal momento che potrebbe avere più effetti negativi che positivi.

In casi in cui la presenza di disclosure è necessaria, un'implicazione che deriva dai risultati della tesi è che l'etichetta andrebbe costruita e presentata cercando di minimizzare il più possibile i suoi effetti

negativi riguardo autenticità, fiducia e credibilità e cercando di non creare reazioni difensive da parte dei consumatori.

Inoltre, i risultati sulla pertinenza context-aware mostrano che l'alta pertinenza contestuale dei contenuti aumenta l'efficacia pubblicitaria perché aumenta utilità e rilevanza percepita.; per i marketer, questo significa che sfruttare le capacità di personalizzazione e contestualizzazione dell'AI in contesti di DOOH può portare a dei grandi risultati di efficacia, permettendo al mezzo di assumere un ruolo centrale all'interno del media mix. Anche in questo caso, se la presenza di disclosure è obbligatoria, è necessario costruirla in modo che i suoi effetti negativi vengano limitati.

In generale, l'implicazione manageriale principale è che, sotto determinate condizioni, il DOOH può essere media molto efficace all'interno delle strategie di comunicazione, soprattutto se implementato con AI.

5.4 Limitazioni della ricerca

Pur offrendo risultati coerenti e informativi, la presente ricerca presenta alcune limitazioni, che sono importanti da sottolineare per interpretare i risultati e orientare le ricerche future.

Una prima limitazione riguarda il campione utilizzato per condurre gli esperimenti. Il campione è stato reclutato online e, quindi, non è un campione probabilistico della popolazione. In più, è un campione composto da relativamente poche unità, il che comporta che l'analisi non è del tutto rappresentativa e i risultati non sono del tutto generalizzabili.

Inoltre, le attitudini verso l'AI e la familiarità con strumenti di AI generativa possono variare molto tra individui, portando a dei risultati diversi; future ricerche potrebbero considerare campioni più robusti per rendere l'analisi più generalizzabile.

Un altro limite riguarda il fatto che l'esposizione agli stimoli è avvenuta in un contesto controllato (questionario online) e non in un reale ambiente urbano con distrazioni, movimento e frammentazione dell'attenzione. Per rendere i risultati più corretti, gli esperimenti dovrebbero essere proposti in un contesto reale.

Un'ulteriore limitazione riguarda la disclosure, costruita come un'etichetta specifica ("Creato con AI"), in una forma standardizzata. Tuttavia, esistono molte forme possibili di disclosure e i risultati, quindi, descrivono l'effetto della disclosure così come implementata e non consentono di concludere che ogni implementazione della disclosure generi gli stessi effetti.

Allo stesso modo, gli stimoli proposti avevano uno specifico formato e a una specifica impostazione creativa. È possibile che in altre categorie di prodotto o con stimoli di diverso formato, gli effetti possano essere diversi.

Un limite riguarda anche il modo in cui è stata misurata la memoria di brand. La brand recognition è stata misurata immediatamente dopo l'esposizione allo stimolo; questa misurazione è coerente con l'obiettivo di misurare effetti immediati post-esposizione, ma non cattura la memoria a distanza di tempo e le forme di memoria più profonde, come la brand recall.

Di conseguenza, i risultati parlano soprattutto di efficacia a breve termine e di impatti diretti sull'encoding immediato.

5.5 Direzioni di ricerca future

A partire dai risultati della ricerca e dalle limitazioni individuate, emergono diverse direzioni di ricerca future che potrebbero generalizzare i risultati e ampliare la ricerca sull'interazione tra AI e DOOH.

Una priorità di ricerca futura è condurre esperimenti sul campo in contesti urbani reali, utilizzando schermi DOOH e misure dell'efficacia più specifiche, come, ad esempio, brand recall a distanza, visite al sito web e engagement sui canali social del brand. Questo permetterebbe di testare se l'origine creativa, la disclosure e la pertinenza context-aware abbiano gli stessi effetti anche in un contesto reale.

Future ricerche dovrebbero, poi, confrontare e testare gli effetti di diverse forme di disclosure all'interno di contenuti DOOH. Questo aiuterebbe a chiarire e generalizzare in quali condizioni la disclosure penalizza realmente l'efficacia pubblicitaria.

Un'ulteriore direzione di ricerca riguarda il testare gli effetti di contenuti DOOH generati con diversi livelli di contributo AI, distinguendo tra contenuto totalmente AI-generated, contenuto co-creato da

AI e umano e contenuto solo ottimizzato con AI. Questo permetterebbe di avere un quadro più completo delle reazioni dei consumatori ai contenuti AI-based.

La letteratura suggerisce che le persone reagiscono diversamente all'utilizzo di AI per la creazione di contenuti in base al fatto che il messaggio creato sia tecnico e funzionale oppure più creativo e emotivo. Questo porta alla possibilità di ricerche future che possano testare se gli effetti negativi della presenza di disclosure in contenuti DOOH siano effettivamente più forti per messaggi valoriali e meno forte per messaggi funzionali.

Infine, potrebbe essere utile studiare gli effetti nel lungo periodo: la disclosure può avere un costo immediato in termini di Aad, ma potrebbe contribuire a costruire fiducia nel brand se integrata in una strategia coerente e più ampia.

5.6 Conclusioni finali

In conclusione, la presente tesi ha analizzato come l'utilizzo di Generative AI per la creazione di contenuti DOOH possa influenzare l'efficacia pubblicitaria del mezzo e, allo stesso tempo, ha studiato quali effetti appaiono quando l'utilizzo di AI diventa visibile ed esplicito per il consumatore. I risultati di questa analisi hanno nel complesso evidenziato che l'AI può essere un abilitatore di efficacia sotto determinate condizioni, ma in altre potrebbe portare a degli effetti negativi.

Da un lato, quando l'origine creativa non è conosciuta, la creatività AI-generated può risultare comparabile a quella umana in termini di efficacia spiegata attraverso l'atteggiamento verso l'annuncio e la memoria di marca e questo suggerisce che l'AI possa essere utilizzata come leva creativa strategica nel DOOH, senza necessariamente compromettere l'efficacia del messaggio.

Dall'altro lato, però, quando l'utilizzo di AI è esplicitato, emergono dei costi significativi, come la diminuzione di autenticità, fiducia e credibilità, l'aumentano della persuasion knowledge e della reactance e una riduzione dell'efficacia complessiva dell'annuncio.

Infine, la pertinenza context-aware si conferma una leva molto importante nel DOOH. Quando il messaggio è percepito come utile e rilevante rispetto al contesto, aumenta sia l'atteggiamento verso l'annuncio sia la memoria di marca; tuttavia, anche qui, la disclosure può attenuare parte dei benefici.

Nel complesso, la tesi chiarisce che l'applicazione dell'Intelligenza Artificiale al DOOH può essere efficace, ma deve essere attentamente costruita e studiata affinché i suoi effetti siano realmente capaci di ridare valore e centralità al mezzo.

BIBLIOGRAFIA

Aguirre, E., Mahr, D., Grewal, D., de Ruyter, K., & Wetzels, M. (2015). Unraveling the personalization paradox: The effect of information collection and trust-building strategies on online advertisement effectiveness. *Journal of Retailing*, 91(1), 34–49.

Araujo, T. (2018). Living up to the chatbot hype: The influence of anthropomorphic design cues and communicative agency framing on conversational agent and company perceptions. *Computers in Human Behavior*, 85, 183–189.

Baek, T. H., Kim, J., & Kim, J. H. (2024). Effect of disclosing AI-generated content on prosocial advertising evaluation. *International Journal of Advertising*.

Bleier, A., & Eisenbeiss, M. (2015). Personalized online advertising effectiveness: The interplay of what, when, and where. *Marketing Science*, 34(5), 669–688.

Brehm, J. W. (1966). *A theory of psychological reactance*. Academic Press.

Brown, S. P., & Stayman, D. M. (1992). Antecedents and consequences of attitude toward the ad: A meta-analysis. *Journal of Consumer Research*, 19(1), 34–51.

Brüns, J. D., & Meißner, M. (2024). Do you create your content yourself? How disclosing generative AI content creation impacts perceived brand authenticity. *Journal of Retailing and Consumer Services*.

Brüns, J. D., & Meißner, M. (2024). Do you create your content yourself? Using generative artificial intelligence for social media content creation diminishes perceived brand authenticity. *Journal of Retailing and Consumer Services*.

Campbell, M. C., & Kirmani, A. (2000). Consumers' use of persuasion knowledge: The effects of accessibility and cognitive capacity on perceptions of an influence agent. *Journal of Consumer Research*, 27(1), 69–83.

- Castelo, N., Bos, M. W., & Lehmann, D. R. (2019). Task-dependent algorithm aversion. *Journal of Marketing Research*.
- Chaudhuri, A., & Holbrook, M. B. (2001). The chain of effects from brand trust and brand affect to brand performance: The role of brand loyalty. *Journal of Marketing*, 65(2), 81–93.
- Chintalapati, S., & Pandey, S. (2022). Artificial intelligence in marketing: A systematic literature review and future research agenda. *Journal of Marketing Management*.
- Cohen, J. (1960). A coefficient of agreement for nominal scales. *Educational and Psychological Measurement*, 20(1), 37–46.
- Cronbach, L. J. (1951). Coefficient alpha and the internal structure of tests. *Psychometrika*, 16(3), 297–334.
- Davenport, T. H., Guha, A., Grewal, D., & Bressgott, T. (2020). How artificial intelligence will change the future of marketing. *Journal of the Academy of Marketing Science*, 48(1), 24–42.
- Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319–340.
- Dennis, C., Brakus, J. J., Gupta, S., & Alamanos, E. (2014). The effect of digital signage on shoppers' behavior: The role of the evoked experience. *Journal of Business Research*, 67(11), 2250–2257.
- Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015). Algorithm aversion: People erroneously avoid algorithms after seeing them err. *Journal of Experimental Psychology: General*, 144(1), 114–126.
- Dillard, J. P., & Shen, L. (2005). On the nature of reactance and its role in persuasive health communication. *Communication Monographs*.
- Ducoffe, R. H. (1996). Advertising value and advertising on the Web. *Journal of Advertising Research*, 36(5), 21–35.

- Edwards, S. M., Li, H., & Lee, J.-H. (2002). Forced exposure and psychological reactance: Antecedents and consequences of the perceived intrusiveness of pop-up ads. *Journal of Advertising*, 31(3), 83–95.
- Erdem, T., & Swait, J. (2004). Brand credibility, brand consideration, and choice. *Journal of Consumer Research*, 31(1), 191–198.
- European Commission. (2025). Code of Practice on marking and labelling of AI-generated content. *ArtificialIntelligenceAct.eu*. (n.d.). Article 50: Transparency obligations for providers and deployers.
- Friestad, M., & Wright, P. (1994). The persuasion knowledge model: How people cope with persuasion attempts. *Journal of Consumer Research*, 21(1), 1–31.
- Garaus, M., Wagner, U., & Rainer, R. C. (2021). Emotional targeting using digital signage systems and facial recognition at the point-of-sale. *Journal of Business Research*, 131, 747–762.
- Grewal, D., Saturnino, C. B., Davenport, T., & Guha, A. (2025). How generative AI is shaping the future of marketing. *Journal of the Academy of Marketing Science*, 53(3), 702–722.
- Ha, L. (1996). Advertising clutter in consumer magazines: Dimensions and effects. *Journal of Advertising Research*, 36(4), 76–84.
- Hair, J. F., Black, W. C., Babin, B. J., & Anderson, R. E. (2019). *Multivariate data analysis* (8th ed.). Cengage.
- Hayes, A. F. (2018). *Introduction to mediation, moderation, and conditional process analysis: A regression-based approach* (2nd ed.). Guilford Press.
- Huang, M.-H., & Rust, R. T. (2021). A strategic framework for artificial intelligence in marketing. *Journal of the Academy of Marketing Science*, 49(1), 30–50.
- Huh, J., & Malthouse, E. C. (2020). Advancing computational advertising: Conceptualization of the field and future directions. *Journal of Advertising*, 49(4), 367–376.

- Hühn, A. E., Khan, V.-J., Ketelaar, P., van 't Riet, J., König, R., Rozendaal, E., Batalas, N., & Markopoulos, P. (2017). Does location congruence matter? A field study on the effects of location-based advertising on perceived ad intrusiveness, relevance & value. *Computers in Human Behavior*.
- Keller, K. L. (1993). Conceptualizing, measuring, and managing customer-based brand equity. *Journal of Marketing*, 57(1), 1–22.
- Kirmani, A., & Rao, A. R. (2000). No pain, no gain: A critical review of the literature on signaling unobservable product quality. *Journal of Marketing*, 64(2), 66–79.
- Kline, R. B. (2016). *Principles and practice of structural equation modeling* (4th ed.). Guilford Press.
- Lang, A. (2000). The limited capacity model of mediated message processing. *Journal of Communication*, 50(1), 46–70.
- Lee, H., & Cho, C.-H. (2019). An empirical investigation on the antecedents of consumers' cognitions of and attitudes towards digital signage advertising. *International Journal of Advertising*, 38(1), 97–115.
- Lee, S. Y. (2025). Maximizing reach and efficiency of digital out-of-home advertising: Quantifying the causal impact of digital billboards on consumer behavior (Working paper).
- Li, H., Edwards, S. M., & Lee, J.-H. (2002). Measuring the intrusiveness of advertisements: Scale development and validation. *Journal of Advertising*, 31(2), 37–47.
- Little, R. J. A., & Rubin, D. B. (2002). *Statistical analysis with missing data* (2nd ed.). Wiley.
- Logg, J. M., Minson, J. A., & Moore, D. A. (2019). Algorithm appreciation: People prefer algorithmic to human judgment. *Organizational Behavior and Human Decision Processes*, 151, 90–103.
- MacKenzie, S. B., & Lutz, R. J. (1989). An empirical examination of the structural antecedents of attitude toward the ad in an advertising pretesting context. *Journal of Marketing*, 53(2), 48–65.

- MacKenzie, S. B., Lutz, R. J., & Belch, G. E. (1986). The role of attitude toward the ad as a mediator of advertising effectiveness: A test of competing explanations. *Journal of Marketing Research*, 23(2), 130–143.
- Mariani, M. M., Perez-Vega, R., & Wirtz, J. (2022). AI in marketing, consumer research and psychology: A systematic literature review and research agenda. *Psychology & Marketing*.
- Mitchell, A. A., & Olson, J. C. (1981). Are product attribute beliefs the only mediator of advertising effects on brand attitude? *Journal of Marketing Research*, 18(3), 318–332.
- Morhart, F., Malär, L., Guèvremont, A., Girardin, F., & Grohmann, B. (2015). Brand authenticity: An integrative framework and measurement scale. *Journal of Consumer Psychology*, 25(2), 200–218.
- Napoli, J., Dickinson, S. J., Beverland, M. B., & Farrelly, F. (2014). Measuring consumer-based brand authenticity. *Journal of Business Research*, 67(6), 1090–1098.
- Obermiller, C., & Spangenberg, E. R. (1998). Development of a scale to measure consumer skepticism toward advertising. *Journal of Consumer Psychology*, 7(2), 159–186.
- Ohanian, R. (1990). Construction and validation of a scale to measure celebrity endorsers' perceived expertise, trustworthiness, and attractiveness. *Journal of Advertising*, 19(3), 39–52.
- Petty, R. E., & Cacioppo, J. T. (1986). *Communication and persuasion: Central and peripheral routes to attitude change*. Springer.
- Pieters, R., & Wedel, M. (2004). Attention capture and transfer in advertising: Brand, pictorial, and text-size effects. *Journal of Marketing*, 68(2), 36–50.
- Preacher, K. J., & Hayes, A. F. (2008). Asymptotic and resampling strategies for assessing and comparing indirect effects in multiple mediator models. *Behavior Research Methods*, 40(3), 879–891.
- Puntoni, S., Reczek, R. W., Giesler, M., & Botti, S. (2021). Consumers and artificial intelligence: An experiential perspective. *Journal of Marketing*, 85(1), 131–151.

Spence, M. (1973). Job market signaling. *The Quarterly Journal of Economics*, 87(3), 355–374.

Stalder, U. (2011). Digital out-of-home media: Means and effects of digital media in public space. In J. Müller, F. Alt, & D. Michelis (Eds.), *Pervasive Advertising* (pp. 31–56). Springer.

van de Sanden, S., et al. (2020). How do consumers process digital display ads in-store? The effect of location, content, and goal relevance. *Journal of Business Research*.

van Meurs, L. V., & Aristoff, M. (2009). Split-second recognition: What makes outdoor advertising work? *Journal of Advertising Research*, 49(1), 82–92.

Wilson, R. T. (2023). Out-of-Home Advertising: A systematic review and research agenda. *Journal of Advertising*.

Wilson, R. T., Baack, D. W., Till, B. D. (2015). Creativity, attention and the memory for brands: An outdoor advertising field study. *International Journal of Advertising*.

Zajonc, R. B. (1968). Attitudinal effects of mere exposure. *Journal of Personality and Social Psychology*, 9(2, Pt. 2), 1–27.

APPENDICE A – Output SPSS (Studio 1)

Descrizione campione - età

Descriptives

		Statistic	Std. Error
Età	Mean	32,03	,787
	95% Confidence Interval for Mean	Lower Bound	30,48
		Upper Bound	33,58
	5% Trimmed Mean	31,05	
	Median	27,00	
	Variance	155,983	
	Std. Deviation	12,489	
	Minimum	19	
	Maximum	70	
	Range	51	
	Interquartile Range	16	
	Skewness	1,194	,153
	Kurtosis	,398	,306

Descrizione campione – genere e istruzione

Genere

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Femmina	120	47,6	47,6	47,6
	Maschio	132	52,4	52,4	100,0
	Total	252	100,0	100,0	

Titolo_studio

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Diploma (scuola secondario di II grado)	104	41,3	41,3	41,3
	Laurea magistrale	68	27,0	27,0	68,3
	Laurea triennale	60	23,8	23,8	92,1
	Licenza media (scuola secondaria di I grado)	8	3,2	3,2	95,2
	Mater / Dottorato	12	4,8	4,8	100,0
	Total	252	100,0	100,0	

Variabili descrittive - indici

Descriptive Statistics

	N	Minimum	Maximum	Mean	Std. Deviation
Aad	252	1,50	7,00	5,0873	1,52153
Auth	252	1,00	6,75	4,7063	1,78615
Trust	252	1,00	6,67	4,5873	1,66384
Cred	252	1,00	7,00	5,0106	1,74716
PKM	252	1,67	7,00	3,8836	1,50685
Reac	252	1,00	7,00	2,3452	1,91889
Valid N (listwise)	252				

Frequenza – brand recognition

Frequency Table

Recognition

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	COVER	216	85,7	85,7	85,7
	Non ricordo	36	14,3	14,3	100,0
	Total	252	100,0	100,0	

0 = non corretto, 1 = corretto

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	,00	36	14,3	14,3	14,3
	1,00	216	85,7	85,7	100,0
	Total	252	100,0	100,0	

Reliability Statistics

Cronbach's Alpha	Cronbach's Alpha Based on Standardized Items	N of Items
,911	,920	4

Item Statistics

	Mean	Std. Deviation	N
Aad_1	5,21	1,385	252
Aad_2	5,27	1,900	252
Aad_3	4,75	1,405	252
Aad_4	5,13	2,055	252

Inter-Item Correlation Matrix

	Aad_1	Aad_2	Aad_3	Aad_4
Aad_1	1,000	,669	,764	,808
Aad_2	,669	1,000	,581	,857
Aad_3	,764	,581	1,000	,767
Aad_4	,808	,857	,767	1,000

Item-Total Statistics

	Scale Mean if Item Deleted	Scale Variance if Item Deleted	Corrected Item-Total Correlation	Squared Multiple Correlation	Cronbach's Alpha if Item Deleted
Aad_1	15,14	24,027	,817	,704	,888
Aad_2	15,08	20,121	,781	,748	,895
Aad_3	15,60	24,559	,754	,665	,905
Aad_4	15,22	17,098	,925	,864	,842

Scale Statistics

Mean	Variance	Std. Deviation	N of Items
20,35	37,041	6,086	4

Analisi di affidabilità - Aad

Reliability Statistics

Cronbach's Alpha	Cronbach's Alpha Based on Standardized Items	N of Items
,972	,973	4

Item Statistics

	Mean	Std. Deviation	N
Auth_1	4,70	1,835	252
Auth_2	4,78	1,885	252
Auth_3	4,27	2,006	252
Auth_4	5,08	1,697	252

Inter-Item Correlation Matrix

	Auth_1	Auth_2	Auth_3	Auth_4
Auth_1	1,000	,953	,927	,878
Auth_2	,953	1,000	,897	,883
Auth_3	,927	,897	1,000	,870
Auth_4	,878	,883	,870	1,000

Item-Total Statistics

	Scale Mean if Item Deleted	Scale Variance if Item Deleted	Corrected Item-Total Correlation	Squared Multiple Correlation	Cronbach's Alpha if Item Deleted
Auth_1	14,13	28,797	,959	,935	,956
Auth_2	14,05	28,476	,945	,917	,959
Auth_3	14,56	27,499	,928	,873	,966
Auth_4	13,75	31,138	,899	,811	,973

Scale Statistics

Mean	Variance	Std. Deviation	N of Items
18,83	51,045	7,145	4

Analisi di affidabilità - Autenticità

Reliability Statistics

Cronbach's Alpha	Cronbach's Alpha Based on Standardized Items	N of Items
,926	,939	3

Item Statistics

	Mean	Std. Deviation	N
Trust_1	4,37	1,497	252
Trust_2	4,56	1,644	252
Trust_3	4,84	2,144	252

Inter-Item Correlation Matrix

	Trust_1	Trust_2	Trust_3
Trust_1	1,000	,869	,738
Trust_2	,869	1,000	,902
Trust_3	,738	,902	1,000

Item-Total Statistics

	Scale Mean if Item Deleted	Scale Variance if Item Deleted	Corrected Item-Total Correlation	Squared Multiple Correlation	Cronbach's Alpha if Item Deleted
Trust_1	9,40	13,659	,815	,767	,931
Trust_2	9,21	11,575	,951	,905	,819
Trust_3	8,92	9,221	,852	,823	,928

Scale Statistics

Mean	Variance	Std. Deviation	N of Items
13,76	24,915	4,992	3

Analisi di affidabilità - Fiducia

Reliability Statistics

Cronbach's Alpha	Cronbach's Alpha Based on Standardized Items	N of Items
,965	,968	3

Item Statistics

	Mean	Std. Deviation	N
Cred_1	4,95	1,842	252
Cred_2	5,14	1,615	252
Cred_3	4,94	1,946	252

Inter-Item Correlation Matrix

	Cred_1	Cred_2	Cred_3
Cred_1	1,000	,897	,924
Cred_2	,897	1,000	,910
Cred_3	,924	,910	1,000

Item-Total Statistics

	Scale Mean if Item Deleted	Scale Variance if Item Deleted	Corrected Item-Total Correlation	Squared Multiple Correlation	Cronbach's Alpha if Item Deleted
Cred_1	10,08	12,121	,933	,872	,944
Cred_2	9,89	13,804	,921	,850	,960
Cred_3	10,10	11,338	,942	,887	,941

Scale Statistics

Mean	Variance	Std. Deviation	N of Items
15,03	27,473	5,241	3

Analisi di affidabilità - Credibilità

Reliability Statistics

Cronbach's Alpha	Cronbach's Alpha Based on Standardized Items	N of Items
,907	,913	3

Item Statistics

	Mean	Std. Deviation	N
PKM_1	4,17	1,401	252
PKM_2	3,43	1,853	252
PKM_3	4,05	1,640	252

Inter-Item Correlation Matrix

	PKM_1	PKM_2	PKM_3
PKM_1	1,000	,690	,829
PKM_2	,690	1,000	,816
PKM_3	,829	,816	1,000

Item-Total Statistics

	Scale Mean if Item Deleted	Scale Variance if Item Deleted	Corrected Item-Total Correlation	Squared Multiple Correlation	Cronbach's Alpha if Item Deleted
PKM_1	7,48	11,087	,792	,687	,895
PKM_2	8,22	8,460	,792	,667	,900
PKM_3	7,60	8,973	,892	,801	,798

Scale Statistics

Mean	Variance	Std. Deviation	N of Items
11,65	20,435	4,521	3

Analisi di affidabilità - PKM

Reliability Statistics

Cronbach's Alpha	Cronbach's Alpha Based on Standardized Items	N of Items
,975	,980	4

Item Statistics

	Mean	Std. Deviation	N
Reac_1	2,43	2,041	252
Reac_2	2,35	1,924	252
Reac_3	2,60	2,283	252
Reac_4	2,00	1,656	252

Inter-Item Correlation Matrix

	Reac_1	Reac_2	Reac_3	Reac_4
Reac_1	1,000	,964	,946	,896
Reac_2	,964	1,000	,920	,950
Reac_3	,946	,920	1,000	,860
Reac_4	,896	,950	,860	1,000

Item-Total Statistics

	Scale Mean if Item Deleted	Scale Variance if Item Deleted	Corrected Item-Total Correlation	Squared Multiple Correlation	Cronbach's Alpha if Item Deleted
Reac_1	6,95	32,301	,968	,954	,959
Reac_2	7,03	33,497	,975	,968	,957
Reac_3	6,78	30,293	,931	,896	,975
Reac_4	7,38	37,559	,917	,909	,978

Scale Statistics

Mean	Variance	Std. Deviation	N of Items
9,38	58,914	7,676	4

Analisi di affidabilità – Reactance

Test H1a

Group Statistics

	Cond_S1	N	Mean	Std. Deviation	Std. Error Mean
Aad	1,00	59	5,7669	,64971	,08458
	2,00	64	5,9531	,57196	,07149

Independent Samples Test

		t-test for Equality of Means						95% Confidence Interval of the Difference	
		t	df	Significance One-Sided p	Significance Two-Sided p	Mean Difference	Std. Error Difference	Lower	Upper
Aad	Equal variances assumed	-1,690	121	,047	,094	-,18618	,11018	-,40430	,03195
	Equal variances not assumed	-1,681	115,979	,048	,095	-,18618	,11075	-,40554	,03318

Test H1b

Cond_S1 * Rec_bin Crosstabulation

		Rec_bin		Total
		,00	1,00	
Cond_S1	1,00	Count	2	57
		% within Cond_S1	3,4%	96,6%
	2,00	Count	8	56
		% within Cond_S1	12,5%	87,5%
Total	Count	10	113	123
	% within Cond_S1	8,1%	91,9%	100,0%

Chi-Square Tests

	Value	df	Asymptotic Significance (2-sided)	Exact Sig. (2-sided)	Exact Sig. (1-sided)
Pearson Chi-Square	3,411 ^a	1	,065		
Continuity Correction ^b	2,301	1	,129		
Likelihood Ratio	3,660	1	,056		
Fisher's Exact Test				,098	,062
Linear-by-Linear Association	3,384	1	,066		
N of Valid Cases	123				

Test H2a

Group Statistics

	Cond	N	Mean	Std. Deviation	Std. Error Mean
Aad	C2	64	5,9531	,57196	,07149
	C3	60	2,7500	1,13496	,14652

Independent Samples Test

		t-test for Equality of Means						95% Confidence Interval of the Difference	
		t	df	Significance One-Sided p	Significance Two-Sided p	Mean Difference	Std. Error Difference	Lower	Upper
Aad	Equal variances assumed	20,031	122	<,001	<,001	3,20313	,15991	2,88657	3,51968
	Equal variances not assumed	19,647	85,880	<,001	<,001	3,20313	,16303	2,87902	3,52723

Test H2b

Cond * brand_rec_bin Crosstabulation

		brand_rec_bin		Total
		non corretto	corretto	
Cond	C2	Count	8	56
		% within Cond	12,5%	87,5%
	C3	Count	28	32
		% within Cond	46,7%	53,3%
Total	Count	36	88	124
	% within Cond	29,0%	71,0%	100,0%

Chi-Square Tests

	Value	df	Asymptotic Significance (2-sided)	Exact Sig. (2-sided)	Exact Sig. (1-sided)
Pearson Chi-Square	17,546 ^a	1	<,001		
Continuity Correction ^b	15,927	1	<,001		
Likelihood Ratio	18,268	1	<,001		
Fisher's Exact Test				<,001	<,001
Linear-by-Linear Association	17,404	1	<,001		
N of Valid Cases	124				

Test H3

Group Statistics

	Cond	N	Mean	Std. Deviation	Std. Error Mean
Auth	C2	64	5,4375	,42258	,05282
	C3	60	1,7167	,54358	,07018

Independent Samples Test

		t-test for Equality of Means						95% Confidence Interval of the Difference	
		t	df	Significance		Mean Difference	Std. Error Difference	Lower	Upper
Auth	Equal variances assumed	42,703	122	<,001	<,001	3,72083	,08713	3,54835	3,89332
	Equal variances not assumed	42,362	111,327	<,001	<,001	3,72083	,08783	3,54679	3,89488

Group Statistics

	Cond	N	Mean	Std. Deviation	Std. Error Mean
Trust	C2	64	5,7500	,53781	,06723
	C3	60	1,8444	,51882	,06698

Independent Samples Test

		t-test for Equality of Means						95% Confidence Interval of the Difference	
		t	df	Significance		Mean Difference	Std. Error Difference	Lower	Upper
Trust	Equal variances assumed	41,107	122	<,001	<,001	3,90556	,09501	3,71748	4,09364
	Equal variances not assumed	41,155	121,897	<,001	<,001	3,90556	,09490	3,71769	4,09342

Group Statistics

	Cond	N	Mean	Std. Deviation	Std. Error Mean
Cred	C2	64	6,0417	,61578	,07697
	C3	60	2,0889	,60713	,07838

Independent Samples Test

		t-test for Equality of Means						95% Confidence Interval of the Difference	
		t	df	Significance		Mean Difference	Std. Error Difference	Lower	Upper
Cred	Equal variances assumed	35,965	122	<,001	<,001	3,95278	,10991	3,73521	4,17035
	Equal variances not assumed	35,981	121,684	<,001	<,001	3,95278	,10986	3,73530	4,17025

Test H4

***** TOTAL, DIRECT, AND INDIRECT EFFECTS OF X ON Y *****

Total effect of X on Y

Effect	se	t	p	LLCI	ULCI
-3,2031	,1599	-20,0308	,0000	-3,5197	-2,8866

Direct effect of X on Y

Effect	se	t	p	LLCI	ULCI
-,4709	,5877	-,8013	,4245	-1,6343	,6925

Indirect effect(s) of X on Y:

	Effect	BootSE	BootLLCI	BootULCI
Auth	-2,7323	,3841	-3,5386	-2,0429

***** TOTAL, DIRECT, AND INDIRECT EFFECTS OF X ON Y *****

Total effect of X on Y

Effect	se	t	p	LLCI	ULCI
-3,2031	,1599	-20,0308	,0000	-3,5197	-2,8866

Direct effect of X on Y

Effect	se	t	p	LLCI	ULCI
,5665	,5069	1,1175	,2660	-,4371	1,5701

Indirect effect(s) of X on Y:

	Effect	BootSE	BootLLCI	BootULCI
Trust	-3,7696	,4857	-4,7793	-2,8379

***** TOTAL, DIRECT, AND INDIRECT EFFECTS OF X ON Y *****

Total effect of X on Y

Effect	se	t	p	LLCI	ULCI
-3,2031	,1599	-20,0308	,0000	-3,5197	-2,8866

Direct effect of X on Y

Effect	se	t	p	LLCI	ULCI
-2,1818	,5382	-4,0536	,0001	-3,2474	-1,1162

Indirect effect(s) of X on Y:

	Effect	BootSE	BootLLCI	BootULCI
Cred	-1,0213	,4489	-1,7978	-,0230

Test H5

Group Statistics

	Cond	N	Mean	Std. Deviation	Std. Error Mean
PKM	C2	64	3,1875	,82268	,10284
	C3	60	5,9778	,79987	,10326
Reac	C2	64	1,2813	,40703	,05088
	C3	60	5,5500	,89821	,11596

Independent Samples Test

		t-test for Equality of Means						95% Confidence Interval of the Difference	
		t	df	Significance		Mean Difference	Std. Error Difference	Lower	Upper
				One-Sided p	Two-Sided p				
PKM	Equal variances assumed	-19,129	122	<,001	<,001	-2,79028	,14587	-3,07904	-2,50152
	Equal variances not assumed	-19,146	121,834	<,001	<,001	-2,79028	,14573	-3,07878	-2,50178
Reac	Equal variances assumed	-34,441	122	<,001	<,001	-4,26875	,12394	-4,51411	-4,02339
	Equal variances not assumed	-33,711	81,089	<,001	<,001	-4,26875	,12663	-4,52070	-4,01680

APPENDICE B – Output SPSS (Studio 2)

Descrizione campione – età

Descriptives

			Statistic	Std. Error
Età	Mean		34,21	,965
	95% Confidence Interval for Mean	Lower Bound	32,31	
		Upper Bound	36,12	
	5% Trimmed Mean		32,86	
	Median		28,00	
	Variance		208,743	
	Std. Deviation		14,448	
	Minimum		21	
	Maximum		73	
	Range		52	
	Interquartile Range		17	
	Skewness		1,435	,163
	Kurtosis		1,055	,324

Descrizione campione – genere e istruzione

Frequency Table

Genere

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Femmina	120	53,6	53,6	53,6
	Maschio	104	46,4	46,4	100,0
	Total	224	100,0	100,0	

Istruzione

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Diploma (scuola secondario di II grado)	80	35,7	35,7	35,7
	Laurea magistrale	48	21,4	21,4	57,1
	Laurea triennale	72	32,1	32,1	89,3
	Licenza media (scuola secondaria di I grado)	8	3,6	3,6	92,9
	Mater / Dottorato	16	7,1	7,1	100,0
	Total	224	100,0	100,0	

Variabili descrittive – indici

Descriptive Statistics

	N	Minimum	Maximum	Mean	Std. Deviation
Aad	224	1,00	7,00	3,4018	1,74281
Util_rel	224	1,00	7,00	3,7232	2,65534
Valid N (listwise)	224				

Frequenze – Brand recognition

Rec_bin

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	,00	112	50,0	50,0	50,0
	1,00	112	50,0	50,0	100,0
	Total	224	100,0	100,0	

Frequenze – Condizioni

Statistics

Cond_S2

N	Valid	224
	Missing	0

Cond_S2

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	1,00	56	25,0	25,0	25,0
	2,00	48	21,4	21,4	46,4
	3,00	64	28,6	28,6	75,0
	4,00	56	25,0	25,0	100,0
	Total	224	100,0	100,0	

Analisi di affidabilità – Aad

Reliability Statistics

Cronbach's Alpha	Cronbach's Alpha Based on Standardized Items	N of Items
,960	,967	4

Item Statistics

	Mean	Std. Deviation	N
Aad_1	3,46	1,806	224
Aad_2	3,39	1,804	224
Aad_3	3,75	1,482	224
Aad_4	3,00	2,209	224

Inter-Item Correlation Matrix

	Aad_1	Aad_2	Aad_3	Aad_4
Aad_1	1,000	,924	,874	,890
Aad_2	,924	1,000	,842	,900
Aad_3	,874	,842	1,000	,844
Aad_4	,890	,900	,844	1,000

Item-Total Statistics

	Scale Mean if Item Deleted	Scale Variance if Item Deleted	Corrected Item-Total Correlation	Squared Multiple Correlation	Cronbach's Alpha if Item Deleted
Aad_1	10,14	27,531	,939	,891	,937
Aad_2	10,21	27,649	,933	,883	,939
Aad_3	9,86	31,693	,881	,786	,961
Aad_4	10,61	23,917	,917	,844	,954

Scale Statistics

Mean	Variance	Std. Deviation	N of Items
13,61	48,598	6,971	4

Analisi di affidabilità – Utilità/rilevanza

Reliability Statistics

Cronbach's Alpha	Cronbach's Alpha Based on Standardized Items	N of Items
,997	,997	4

Item Statistics

	Mean	Std. Deviation	N
Utility_1	3,75	2,647	224
Utility_2	3,64	2,600	224
Utility_3	3,75	2,714	224
Utility_4	3,75	2,714	224

Inter-Item Correlation Matrix

	Utility_1	Utility_2	Utility_3	Utility_4
Utility_1	1,000	,983	,990	,985
Utility_2	,983	1,000	,984	,984
Utility_3	,990	,984	1,000	,995
Utility_4	,985	,984	,995	1,000

Item-Total Statistics

	Scale Mean if Item Deleted	Scale Variance if Item Deleted	Corrected Item-Total Correlation	Squared Multiple Correlation	Cronbach's Alpha if Item Deleted
Utility_1	11,14	63,908	,990	,983	,996
Utility_2	11,25	64,762	,987	,974	,997
Utility_3	11,14	62,688	,995	,994	,994
Utility_4	11,14	62,760	,993	,991	,995

Scale Statistics

Mean	Variance	Std. Deviation	N of Items
14,89	112,814	10,621	4

Test H6a

Group Statistics

	Cond_S2	N	Mean	Std. Deviation	Std. Error Mean
Aad	1,00	56	2,9286	,48080	,06425
	2,00	48	6,5000	,38592	,05570

Independent Samples Test

		t-test for Equality of Means						95% Confidence Interval of the Difference	
		t	df	Significance One-Sided p	Significance Two-Sided p	Mean Difference	Std. Error Difference	Lower	Upper
Aad	Equal variances assumed	-41,300	102	<,001	<,001	-3,57143	,08648	-3,74295	-3,39991
	Equal variances not assumed	-42,000	101,590	<,001	<,001	-3,57143	,08503	-3,74010	-3,40276

Test H6b

Cond_S2 * Rec_bin Crosstabulation

			Rec_bin		Total
			,00	1,00	
Cond_S2	1,00	Count	24	32	56
		% within Cond_S2	42,9%	57,1%	100,0%
	2,00	Count	0	48	48
		% within Cond_S2	0,0%	100,0%	100,0%
Total	Count		24	80	104
	% within Cond_S2		23,1%	76,9%	100,0%

Chi-Square Tests

	Value	df	Asymptotic Significance (2-sided)	Exact Sig. (2-sided)	Exact Sig. (1-sided)
Pearson Chi-Square	26,743 ^a	1	<,001		
Continuity Correction ^b	24,383	1	<,001		
Likelihood Ratio	35,877	1	<,001		
Fisher's Exact Test				<,001	<,001
Linear-by-Linear Association	26,486	1	<,001		
N of Valid Cases	104				

Test H7

Group Statistics

	Cond_S2	N	Mean	Std. Deviation	Std. Error Mean
Util_rel	1,00	56	1,5000	,44721	,05976
	2,00	48	6,7083	,26962	,03892

Independent Samples Test

		t-test for Equality of Means						95% Confidence Interval of the Difference	
		t	df	Significance		Mean Difference	Std. Error Difference	Lower	Upper
Util_rel	Equal variances assumed	-70,431	102	<,001	<,001	-5,20833	,07395	-5,35501	-5,06166
	Equal variances not assumed	-73,032	92,146	<,001	<,001	-5,20833	,07132	-5,34997	-5,06670

Test H8a

***** TOTAL, DIRECT, AND INDIRECT EFFECTS OF X ON Y *****					
Total effect of X on Y					
Effect	se	t	p	LLCI	ULCI
2,2295	,1799	12,3912	,0000	1,8749	2,5841
Direct effect of X on Y					
Effect	se	t	p	LLCI	ULCI
-3,4833	1,1687	-2,9806	,0032	-5,7864	-1,1801
Indirect effect(s) of X on Y:					
	Effect	BootSE	BootLLCI	BootULCI	
Util_rel	5,7128	1,0252	3,5963	7,6373	

Test H8b

***** DIRECT AND INDIRECT EFFECTS OF X ON Y *****					
Direct effect of X on Y					
Effect	se	Z	p	LLCI	ULCI
-15,1514	2,4776	-6,1154	,0000	-20,0074	-10,2954
Indirect effect(s) of X on Y:					
	Effect	BootSE	BootLLCI	BootULCI	
Util_rel	18,1029	2,7474	12,6276	23,4816	

Test H9a

Group Statistics

	Cond_S2	N	Mean	Std. Deviation	Std. Error Mean
Aad	2,00	48	6,5000	,38592	,05570
	4,00	56	2,9643	,28376	,03792

Independent Samples Test

		t-test for Equality of Means						95% Confidence Interval of the Difference	
		t	df	Significance One-Sided p	Significance Two-Sided p	Mean Difference	Std. Error Difference	Lower	Upper
Aad	Equal variances assumed	53,701	102	<,001	<,001	3,53571	,06584	3,40512	3,66631
	Equal variances not assumed	52,471	85,046	<,001	<,001	3,53571	,06738	3,40174	3,66969

Test H9b

Cond_S2 * Rec_bin Crosstabulation

		Rec_bin		Total
		,00	1,00	
Cond_S2	2,00	Count	0	48
		% within Cond_S2	0,0%	100,0%
	4,00	Count	24	32
		% within Cond_S2	42,9%	57,1%
Total	Count	24	80	104
	% within Cond_S2	23,1%	76,9%	100,0%

Chi-Square Tests

	Value	df	Asymptotic Significance (2-sided)	Exact Sig. (2-sided)	Exact Sig. (1-sided)
Pearson Chi-Square	26,743 ^a	1	<,001		
Continuity Correction ^b	24,383	1	<,001		
Likelihood Ratio	35,877	1	<,001		
Fisher's Exact Test				<,001	<,001
Linear-by-Linear Association	26,486	1	<,001		
N of Valid Cases	104				