



“RICERCA INTEGRATA”

NOTE INFORMATIVE

Audicom Srl, con sede legale in Milano, Via Larga 13, C.F., n. iscrizione. Reg. Imp. e P.I. 12829780969, è la società risultante dalla fusione per unione di Audiweb Srl e Audipress Srl con data di efficacia 1° marzo 2023.

Sommario

1. SOCIETA'	5
1.1. Missione	5
1.2. Governance	5
1.3. Dati anagrafici	5
2. INFORMAZIONI DI BASE	6
Sistemi di Rilevazione	6
Servizi di Distribuzione	6
Dati e modalità di distribuzione	7
2.1. Fornitori per i sistemi di rilevazione	8
2.2. Accesso ai servizi e tariffe.....	8
2.3. Metriche	17
3. RICERCA DI BASE	18
3.1 Dati anagrafici della Società di ricerca	18
3.2 Obiettivo della rilevazione	18
3.3 Oggetto della ricerca	18
3.4 Metodologia.....	19
3.4.1 Consistenza del campione oggetto di indagine	19
3.4.2 Modalità di rilevazione ed eventuale margine di errore	19
3.4.3 Universo di riferimento.....	21
3.4.4 Metodo di campionamento	22
3.4.5 Ponderazione ed espansione dei risultati	22
3.5 Periodo di rilevazione.....	24
3.6 Indirizzo web di pubblicazione della metodologia.....	24
4. AUDICOM DIGITAL PANEL	25
4.1. Dati anagrafici delle Società di ricerca	25
4.2. Obiettivo della rilevazione	25
4.3. Metodologia.....	25
4.3.1. Software Meter per la rilevazione PC (Desktop e Laptop)	25
4.3.2. Raccolta dei dati panel.....	26
4.3.3. Software Meter per la rilevazione Mobile (Smartphone e tablet).....	28

4.4. Campione oggetto di indagine	28
4.4.1. Single digital recruitment.....	28
4.4.2. Reclutamento dei Minori per la rilevazione mobile	29
4.5. Ponderazione ed eleggibilità per i dati panel PC e Mobile.....	29
4.5.1. Processo di ponderazione per la componente Computer.....	29
4.5.2. Processo di ponderazione per la componente Mobile	32
4.5.3. Criteri di eleggibilità	32
4.5.4. Controlli antifrode sui panelisti.....	33
4.6. Utilizzo del campione convenience digitale nella ricerca stampa CAWI	34
4.7. Periodo di rilevazione	34
4.8. Indirizzo web di pubblicazione della metodologia.....	34
5. AUDICOM CENSUS.....	35
5.1. Dati anagrafici della Società di ricerca	35
5.2. Obiettivo della rilevazione	35
5.2.1. SDK Text – caratteristiche e funzionamento	36
5.2.2. SDK Text– certificazione	38
5.2.3. Attività di controllo a cura di terze parti.....	39
5.2.4. SDK Text – processo di raccolta dati	39
5.2.4.1. Filtro per il traffico non umano (Anti-Fraud).....	40
5.3. “Audicom Daily/Weekly Digital”	40
5.3.1. L’uso di DCR ID per la misurazione dell’audience.....	41
5.3.2. ID Graph.....	41
5.3.2.1. Matrice di correzione delle demografiche	42
5.3.2.2. Matrice della condivisione	42
5.3.2.3. Fattori di Correzione della Non-Coverage	43
5.3.3. Deduplicazione	43
5.3.4. SDK VIDEO (a cura di Auditel S.r.l.).....	45
5.3.4.1. Rilevazione Censuaria.....	45
5.3.4.2. Elaborazione dei dati e fornitura del dato censuario.....	48
5.3.4.3. Elaborazione Dato Editoriale	49
5.3.4.4. Elaborazione Dato Pubblicitario	49

5.3.4.5.	CUSV – Codice Unico Spot Video.....	50
5.3.5.	Integrazione dei Dati Censuari Video (a cura di The Nielsen Media Italy S.r.L.) ...	51
5.3.6.	Distribuzione Dati “Audicom Daily/Weekly Digital”	52
5.4.	Privacy e Riservatezza.....	53
6.	AUDICOM CTV (a cura di The Nielsen Media Italy S.r.L.).....	54
6.1.	Obiettivo della rilevazione	54
6.2.	Metodologia.....	54
6.3.	Distribuzione del Dato	54
6.4.	Periodo di rilevazione	55
6.5.	Indirizzo web di pubblicazione della metodologia.....	55
7.	AUDICOM DATABASE RESPONDENT LEVEL DIGITAL	56
7.1.	Introduzione	56
7.2.	Metodologia Synthetic Respondent Level (per il dato digitale).....	56
7.2.1.	Processo di Fusione dei panel digitali	57
7.2.2.	Soft Calibration	58
7.2.3.	Algoritmo SRLD.....	59
7.3.	Rilevazione e gestione del traffico dei contenuti distribuiti tramite mobile In-App Browsing	61
8.	AUDICOM PRINT SURVEY	62
8.1.	Quadro sintetico	62
8.1.1.	Universo considerato.....	62
8.1.2.	Campione per il rilevamento dei dati ed epoca del rilevamento	62
8.1.3.	Numero dei comuni di campionamento	62
8.2.	Sintesi metodologica	62
8.2.1.	Premessa	62
8.2.2.	Le stime dei lettori	63
8.2.3.	Campionamento	66
8.2.3.1.	Metodo di campionamento CAPI.....	67
8.2.3.2.	Metodo di campionamento CAWI.....	73
8.2.4.	Rilevamento	74
8.2.5.	Stime e definizioni di lettura.....	79

8.2.6. Ripartizione ed esecuzione delle interviste.....	81
8.3. Elaborazioni	85
8.3.1. Attività di controllo a cura di terze parti.....	96
8.4. Definizione e note per l'uso dei dati	96
9. AUDICOM DATABASE RESPONDENT LEVEL DIGITAL & PRINT.....	105
9.1. Fusione rilevazione digital panel e stampa.....	105
9.2. Metodologia per il collegamento delle due ricerche	105
9.3. Data Fusion	106
10. CATALOGAZIONE DELL'OFFERTA EDITORIALE E PUBBLICITARIA .	107
10.1. Costruzione del MarketView	107
10.2. Traffic Assignment Letter (TAL)	113
10.3. Collegamento testate "Print" con rilevazione "Digital" per distribuzione dati "Digital&Print"	115
11. ALLEGATI.....	116
11.1. Schema del Database Respondent Level Digital.....	116
11.2. Schema Navigation DataSet Print	125
11.3. Audicom Print - Gli intervalli fiduciari delle stime.....	126
11.4. Attività di cross-recruitment tra Audicom Digital Panel e Audicom Print Survey ..	131
11.5. Audicom ADV Census - Tracciato record 5.4	133

1. SOCIETA'

1.1. Missione

Audicom è il Joint Industry Committee nato dalla fusione di Audiweb e Audipress con l'obiettivo di realizzare e offrire al mercato una "Ricerca Integrata" sulla fruizione di contenuti multimediali, editoriali e/o pubblicitari, mediante internet e stampa quotidiana e periodica.

1.2. Governance

Audicom è una società partecipata da Fedoweb, associazione di Editori e Operatori web (24,87%), Fieg Federazione Italiana Editori Giornali (24,87%), UPA Utenti Pubblicità Associati (24,87%), che rappresenta le aziende nazionali e multinazionali che investono in pubblicità e Assap Servizi (24,87%), l'azienda servizi di proprietà di UNA, associazione delle agenzie media operanti in Italia. Si configura quindi come un Joint Industry Committee con la partecipazione delle associazioni di categoria di tutti gli operatori del mercato. Partecipa al capitale di Audicom Srl, in virtù di rapporti di reciprocità, anche Auditel con una quota di minoranza (0,5%), senza tuttavia diritti alla designazione dei Consiglieri di Amministrazione.

La Società è gestita da un Consiglio di Amministrazione affiancato da un Comitato Tecnico che ha funzioni propositive e consultive in relazione all'impostazione delle rilevazioni, delle ricerche e della diffusione dei dati ottenuti.

1.3. Dati anagrafici

Audicom Srl

Sede Legale: Via Larga 13, 20122 Milano

Tel. +39 02 58305820

E-mail: info@audicom.net

www.audicom.net

C.F., n. iscrizione. Reg. Imp. e P.I. 12829780969

Capitale sociale euro 60.303,00 i.v.

Presidente: Marco Travaglia

Amministratore Delegato: Marco Muraglia

2. INFORMAZIONI DI BASE

La “Ricerca Integrata” Audicom consente la rilevazione completa dei consumi editoriali DIGITAL “video” e “text” distribuiti via smartphone, tablet, computer e CTV e PRINT su quotidiani e periodici. La ricerca consente inoltre di rilevare il dato di audience unica deduplicata collegando l’audience “digital” con il lettorato “print”, dando origine alla rilevazione “Digital&Print”.

La Ricerca Integrata consente inoltre la rilevazione dei dati censuari dei consumi adv video.

Sistemi di Rilevazione

La “Ricerca Integrata” Audicom si basa sui seguenti “Sistemi di Rilevazione”:

- “Ricerca di Base”: rileva le caratteristiche e descrive l’universo degli individui con possibilità di accesso ad internet. È condotta mediante interviste in modalità CAPI, fornita da Auditel ed operata da IPSOS.
- “Audicom Digital Panel”: rileva in modo continuativo ed oggettivo la navigazione online via computer e mobile, è realizzata mediante “On Device Meter” installati su computer, smartphone e tablet di un campione di individui. Rilevazione operata da Nielsen.
- “Audicom Print Survey”: rileva i lettori della stampa quotidiana e periodica mediante interviste realizzate in modalità CAPI e CAWI. Rilevazione operata da IPSOS DOXA.
- “Audicom Census”: consolida la rilevazione panel e consente completa copertura e coerente attribuzione delle audience. Realizzata mediante “SDK Text” (o “SDK Testo”) operato da Nielsen e “SDK Video” fornito da Auditel e operato da comScore. Alimenta il servizio “Audicom ADV Census” rilevando i consumi di «stream views» e di tempo speso dei contenuti video ADV rilevati mediante “SDK Video” e “CUSV” - Codice Unico Spot Video - che consente di assegnare a ciascun contenuto video ADV un codice collegato ad una specifica creatività.
- “Audicom CTV”: rileva la fruizione di contenuti video online tramite Connected TV. Realizzata da Auditel mediante Panel Focal Meter e SDK Video.

Servizi di Distribuzione

I suddetti “Sistemi di Rilevazione” consentono la rilevazione dei consumi di contenuti digital editoriali e pubblicitari distribuiti via computer, mobile e connected TV mediante web browsing, mobile APP e CTV APP per contenuti sia testuali sia video alimentando i seguenti “servizi di distribuzione” dei dati. Consentono inoltre la rilevazione dei consumi print a sé stanti e/o collegati ai consumi digital.

- “Audicom Print”: Indagine di riferimento per la lettura della stampa quotidiana e periodica.

- “Audicom Digital”: Rilevazione dei dati sulla fruizione dei mezzi operanti su Internet, distribuiti via web browsing e APP
 - “Digital Text”: contenuti editoriali “text” fruiti da computer e mobile
 - “Digital Combined”: contenuti editoriali “Text & Video” fruiti da computer e mobile e, per la sola parte “video”, da CTV
- “Audicom Digital & Print”: consente di stimare l’utenza unica deduplicata fra gli utenti internet ed i lettori stampa per i mezzi iscritti alla “Ricerca Integrata”.
- “Audicom ADV Census”: rileva i consumi di «stream views» e di tempo speso dei contenuti video ADV. Il Servizio consente la misurazione dei volumi censuari (LSVs) e del tempo speso (TS) dei contenuti pubblicitari video erogati via computer, smartphone, tablet e CTV collegati al perimetro editoriale dei Publisher Iscritti e, qualora dotati di CUSV, decodificabili per azienda e soggetto.

Dati e modalità di distribuzione

I “Servizi di Distribuzione” dati di Audicom prevedono diversi di tipi di output di dati:

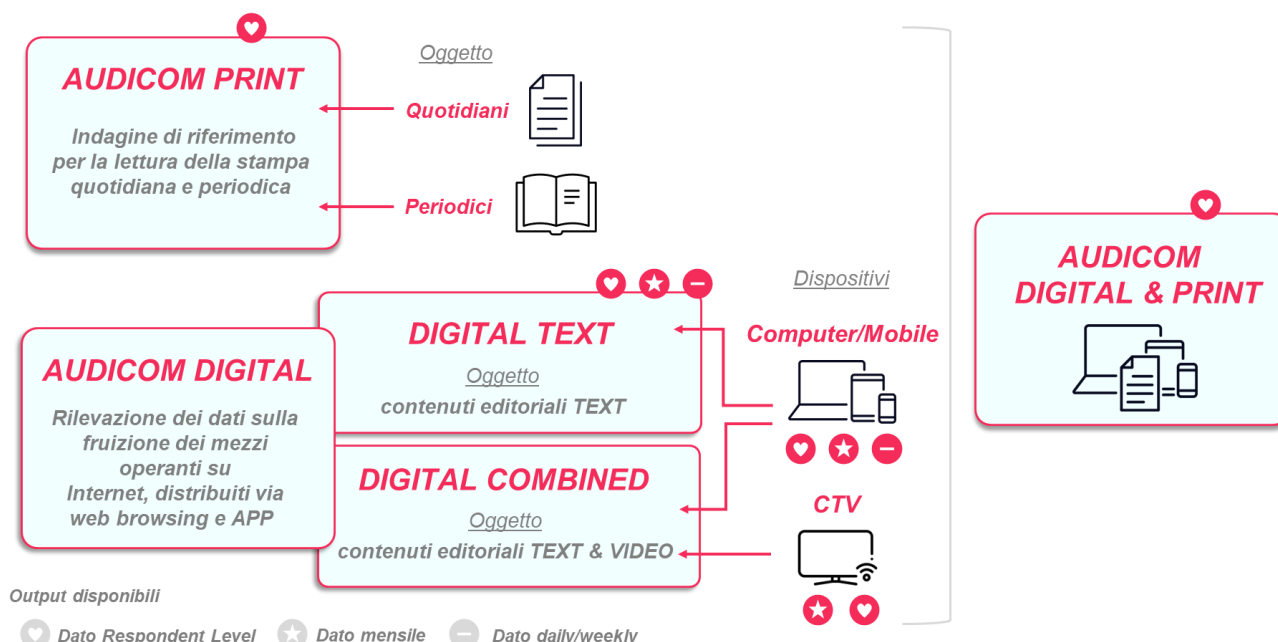
- “Audicom Database Respondent Level Print”: respondent level database distribuito a frequenza quadrimestrale accessibile mediante software predisposti da software house accreditate. Un set di dati analitici derivanti da “Audicom Database Print” viene distribuito mediante il sito della società.
- “Audicom Database Respondent Level Digital”: respondent level database distribuito a frequenza mensile accessibile mediante “Audicom Media View” e software predisposti da software house accreditate.
- “Audicom Daily/Weekly Digital” con dati “giorno” e “settimana” distribuiti a frequenza giornaliera e settimanale per le properties dotate di SDK Text e Video per i consumi da computer e mobile accessibile mediante “Audicom Media View Daily/Weekly”.
- “Audicom Database Respondent Level Digital & Print”: respondent level database distribuito a frequenza mensile accessibile mediante “Audicom Media View” e software predisposti da software house accreditate.
- “Audicom ADV Census”: tracciato record con dati “giorno” distribuito a frequenza giornaliera accessibile mediante software predisposti da software house accreditate. L’output di questo servizio è quindi un file quotidiano (“Tracciato record 5.4” - Allegato 11.5) che per conto di Audicom viene elaborato da Auditel e da Auditel trasmesso alle software house. Il Servizio di Distribuzione “Audicom ADV Census” è basato sul Sistema di Rilevazione “Audicom Census” alimentato per la parte dei contenuti pubblicitari video dall’SDK Video Unico.

2.1. Fornitori per i sistemi di rilevazione

- “Ricerca di Base”: Auditel S.r.l. (ricerca operata da IPSOS)
- “Audicom Digital Panel”: The Nielsen Media Italy S.r.l.
- “Audicom Print Survey”: The Nielsen Media Italy S.r.l. (ricerca operata da IPSOS DOXA)
- “Audicom Census”: The Nielsen Media Italy S.r.l. e Auditel S.r.l.
- “Audicom CTV”: The Nielsen Media Italy S.r.l. e Auditel S.r.l.

2.2. Accesso ai servizi e tariffe

La grafica qui di seguito rappresenta in sintesi i servizi di distribuzione dati sottoscrivibili.



L'accesso ai servizi di distribuzione è offerto a tre diversi profili di sottoscrittori:

Per accedere ai servizi di rilevazione e consultazione dei dati è necessaria la sottoscrizione di un contratto con Audicom. Le tariffe e le opportunità di accesso ai servizi sono diverse a seconda della tipologia del Soggetto che desidera iscriversi:

Per “Audicom Print”:

- Publisher: Editori, Concessionarie
- Requisiti di ammissione alla rilevazione per le testate dei Publisher
 - quotidiani e i settimanali regolarmente registrati ed editi da almeno tre mesi

- mensili e altri mezzi periodici regolarmente registrati ed editi da almeno sei mesi.

Per ogni testata che partecipi per la prima volta all'indagine è dovuta una quota una tantum di euro 5.000,00.

Nel caso di nuove testate o di cambi di periodicità, l'Editore può richiedere ad Audicom la pubblicazione di un dato preliminare (stima dei lettori, con le ripartizioni di base in Uomini, Donne, Responsabili Acquisti), basato sulla cumulazione dei primi due cicli di rilevazione disponibili (l'indagine "Audicom Print" si articola su tre cicli annuali). Il dato preliminare pubblicato non verrà inserito nel nastro dei risultati dell'Indagine.

La disponibilità del primo dato pianificabile avviene dopo l'elaborazione della cumulata di tre cicli di rilevazione.

Per "Audicom Digital":

- **Publisher iscritti.** Editori, Concessionarie, Aggregatori, Service Provider rilevati attraverso Audicom Digital Panel, le cui Property vengono pubblicate in Audicom Database Respondent Level Digital. I Publisher iscritti possono accedere alla rilevazione Audicom Census.
- **Utilizzatori Professionali – Agenzie Media.** Soggetti interessati ai dati per analisi di carattere pubblicitario
- **Altri Utilizzatori**
 - Aziende. Molteplicità di Soggetti che non appartengono ad alcuna delle categorie precedenti e sono interessati ai dati per analisi di carattere scientifico, commerciale, divulgativo (aziende investitrici, società di consulenza, operatori dell'informazione, aziende produttrici di contenuti o servizi non censiti, ecc.).
 - Publisher in CRU. Editori iscritti ad Audicom Database Respondent Level Digital non in virtù di una sottoscrizione diretta (come invece i "Publisher iscritti"), ma all'interno del perimetro di una CRU Concessionaria.

Nuovi sottoscrittori dei servizi di rilevazione dei servizi editoriali. Soggetti che non hanno avuto un contratto annuale attivo con Audicom negli ultimi due anni (2024 e 2025); beneficiano per il 2026 di una riduzione del 30% su tutte le tariffe e per tutti i servizi, valida solo per il primo anno o per il valore a completamento dell'anno solare, in caso di iscrizione ad anno in corso. Tale riduzione si applica anche ai Sottoscrittori che nel 2025 non abbiano aderito ad un intero anno di servizi.

Servizio di distribuzione dati "Audicom Digital Text": modalità di calcolo delle tariffe

Le tariffe sono calcolate in base alle dimensioni di audience e traffico mensile del perimetro delle *Property* iscritte.

Nel caso in cui un Publisher iscritto sia titolare di o rappresenti più *Property*, può riunirle sotto una *Custom Roll-up* (CRU Concessionaria) dove queste saranno rappresentate anche come un'unica entità, seppur rilevate separatamente.

Il Publisher iscritto può definire anche una o più *Custom Roll-up* tematiche (CRU tematiche), ovvero aggregazioni di una serie di Brand e Sub-Brand da esso medesimo iscritti e con tematiche affini che vengono rappresentati unitariamente per scopi commerciali.

La tariffa prevista per il 2026 deriva dalla somma di tre componenti:

- Fissa
- Variabile
- Contributo per l'utilizzo di SDK Text e per il servizio *Audicom Digital Daily/Weekly* (per i Publisher che ne fanno richiesta)

Componente fissa

Per determinare la componente fissa della tariffa, ogni Publisher iscritto viene attribuito ad una classe, in funzione del valore di *Active Reach Total Digital Audience*.

Componente variabile

Per calcolare la componente variabile della tariffa si utilizza una formula che pondera i parametri di *Unique Audience* e *Total Minutes*. Per il 2026 si considera la media dei valori mensili *currency* del primo semestre 2025, come rilevati dal *sistema Audiweb* e disponibili attraverso le piattaforme di consultazione dei dati:

$$\text{Componente variabile} = a * [(total\ minutes\ AwDbMC\ '000/b) + unique\ audience\ AwDbMC\ '000] * comp + a * [((total\ minutes\ TDA - total\ minutes\ AwDbMC\ '000)/b) + (unique\ audience\ TDA - unique\ audience\ AwDbMC\ '000)] * mob$$

Dove:

- AwDbMC = valore da Audiweb Digital Database Monthly, Computer
- TDA = Audiweb Database Monthly, *Total Digital Audience*
- “a” = 5,5 e “b” = 57. Due coefficienti utilizzati per ponderare le componenti *Tempo Speso* e *Unique Audience*
- “comp” = 0,74. Fattore moltiplicativo per la parte di audience Computer
- “mob” = 0,56. Fattore moltiplicativo per la parte di audience incrementale Mobile (audience TDA – audience AwDbMC), uguale al 75% della parte Computer

Solo ed esclusivamente in favore dei Publisher con contratto attivo nel 2025, le tariffe 2026 [Componente fissa + Componente variabile] così ottenute sono soggette ad una riduzione del 4%.

Ai Publisher che utilizzano l'SDK Text – e per i quali viene quindi attivato il servizio *Audicom Daily/Weekly* – viene richiesto un contributo supplementare pari al 10% del valore dell'adesione

delle proprie *Properties* alle rilevazioni e servizi di base (inteso come somma di componente fissa e variabile).

La tariffa per una *Custom Roll-up* (CRU Concessionaria) viene quotata come se questa fosse una *Property* unica: i valori di Tempo Speso di ciascuna entità inserita nella CRU vengono sommati e (nel caso di *Active Reach* e *Unique Audience*) deduplicati.

La sottoscrizione al servizio dà diritto alla pubblicazione dei dati delle *Property* in *Audicom Database Respondent Level Digital* e qualora il Publisher Iscritto abbia installato SDK Text o Video alla rilevazione *Audicom Census* ed al servizio *Audicom Digital Daily/Weekly*.

Schema di sintesi

	Componente		Servizi base inclusi	Accesso ad Audicom Digital Media View daily, weekly, monthly (euro)
Classi (Active Reach TDA)	Fissa (euro)	Variabile		
< 2%	11.500	non previsto	– Rilevazione Audicom Digital Panel – Rilevazione Audicom Census e servizio Audicom Daily/Weekly	+ 17.000
tra 2 e 5%		Da calcolare in funzione della dimensione del perimetro iscritto		
tra 5 e 10%	28.500		– Pubblicazione dei dati in Audicom Database Respondent Level Digital – Licenza d'uso di Audicom Database Respondent Level Digital	Compreso
fra 10 e 20%	35.000			
fra 20 e 45%	42.000			
fra 45 e 60%	60.000			
Oltre 60%	95.000			

Il valore massimo consentito calcolato come somma di [Componente fissa + Componente variabile] è pari a 346.000 euro.

Il valore massimo consentito calcolato come somma di [Componente fissa + Componente variabile + 10% per utilizzo di SDK Text] è pari a 380.000 euro.

Per l'attivazione di *Custom Roll-Up* tematiche, la tariffa è pari a 5.650 euro annui e prevede la possibilità di realizzarne fino ad un massimo di 20.

Servizio di attribuzione a Sub-Brand tematici dei consumi di contenuti editoriali statici / testuali, fruiti mediante Google AMP

Tariffa annua: 5.000 euro da 1 a 5 Sub-Brand, 1.000 euro ulteriori per ogni Sub-Brand aggiuntivo. Tariffa annua mensilizzata in funzione del mese di erogazione del servizio.

-----oooooooooooooooo-----

Servizio di distribuzione dati “Audicom Digital Combined”: modalità di calcolo delle tariffe

Il servizio di distribuzione dati “Audicom Digital Combined” è destinato ai publisher che abbiano sottoscritto il servizio “Audicom Digital Text”, è alimentato dal sistema di rilevazione “Audicom

Census” SDK Video curato da Auditel. L’adesione al servizio prevede tariffe diverse per publisher che aderiscono:

- esclusivamente ai servizi Audicom (“clienti Audicom”)
- sia ai servizi Audicom sia ai servizi Auditel (“clienti condivisi con Auditel”)

Le tariffe comprendono i costi della misurazione delle audience video (SDK Video unico compreso per i sottoscrittori esclusivi Audicom) e della fusione con il dato «Audicom Digital Text» (per entrambe le tipologie di sottoscrittori)

Il valore complessivo della rilevazione mediante SDK Video unico è attribuito ai soli sottoscrittori del servizio ed è definito convenzionalmente e complessivamente a sistema per il 2026 a 250.000,00 euro max ⁽¹⁾ per 400.000.000 stream views mese con un CAP pari a 20.000 euro.

- ✓ Per i “clienti Audicom”, le tariffe prevedono il calcolo di una percentuale della stima di SVs mensili generate nell’anno precedente rispetto al totale SVs mese Audicom.

Tariffa «Audicom Digital Combined»: (SV Publisher/SVs mese)*250.000, per i nuovi clienti in base alla stima dichiarata con consuntivo a fine anno.

- ✓ Per i clienti condivisi con Auditel, non si applicano i costi della rilevazione SDK Video ma i soli costi di elaborazione del dato pari al 5% del totale «Digital Text».

Tariffa «Audicom Digital Combined»: Tariffa Digital Text*0,05

⁽¹⁾ *NOTA: E’ un valore determinato convenzionalmente al solo fine di elaborare le tariffe.*

-----oooooooooooo-----

Servizio di distribuzione dati “Audicom Digital Combined & CTV”: modalità di calcolo delle tariffe

L’adesione al servizio di distribuzione «Audicom Digital Combined & CTV» è vincolata all’adesione al servizio «Audicom Digital Combined».

La tariffa prevista per il servizio di distribuzione «Audicom Digital Combined & CTV» deriva dalla somma di tre componenti:

- ✓ Valore della tariffa «Audicom Digital Text»⁽¹⁾ * a
- ✓ Valore rilevazione CTV: 9,72% del servizio “Audicom Digital Text”

dove «a» = 0,86. Fattore moltiplicativo per la produzione del dato “Audicom Digital Combined & CTV”.

- ✓ Quota fissa (CAP) pari a 180.000 euro per i Publisher (parent) con dati di active reach «Audicom Digital Combined» mese maggiori o uguali al 25% e di 40.000 euro per i Publisher (parent) con dati di active reach «Audicom Digital Combined» mese inferiori al 25%

NOTE:

(1) nel caso di concessionarie che iscrivono publisher rilevati su CTV insieme a publisher «digital» il valore della componente «Audicom Digital Text» è calcolato esclusivamente sul valore «digital text» dei publisher (parent) che richiedono la rilevazione CTV.

-----oooooooooooooooo-----

Schema sintetico calcolo tariffe servizi Audicom Digital

Per «Digital Text»:

- ✓ Valore fisso, determinato dal livello di Active Reach Total
- ✓ Valore variabile, determinato da una formula che pondera audience e tempo speso Total e incrementale Mobile
- ✓ +10% su fisso e variabile per l'utilizzo di SDK text (opzionale)

Il valore minimo per 12 mesi è 11.500 euro, il valore massimo (CAP) è 380.000 euro

Per «Digital Combined (video incluso)»:

- ✓ Per i già clienti Audicom: valore calcolato in base ai consumi Stream Views. Il valore massimo (CAP) è 20.000 euro;
- ✓ Per i clienti condivisi Audicom/Auditel: Tariffa Digital TEXT*a
dove a = 0,05: Fattore moltiplicativo per la parte Digital

Per «Digital Combined + CTV»:

- ✓ (Valore della tariffa «Audicom Digital Text»*a) + Valore rilevazione CTV
- ✓ a = 0,86: Fattore moltiplicativo per produzione del dato «Audicom Digital Video COMBINED + CTV
- ✓ Valore rilevazione CTV = 9,72 % del servizio «Audicom Digital Text»

-----oooooooooooooooo-----

Servizio di distribuzione dati «Audicom Print»: modalità di calcolo delle tariffe

Accesso ai servizi e tariffe

Publisher Iscritti

Le tariffe sono calcolate in base alle dimensioni di readership delle testate e al perimetro delle testate iscritte.

L'ammontare complessivo dell'addebito ai Publisher è deciso ed approvato dal Consiglio di Amministrazione di Audicom che, di anno in anno, provvede a comunicare ai Publisher l'ammontare complessivo richiesto. Tale importo potrà variare in aumento fino ad un massimo del 10% su motivata decisione del Consiglio di Amministrazione di Audicom.

La tariffa prevista per il 2026 prevede i seguenti criteri di suddivisione dei costi per le testate partecipanti:

- Quotidiani/Supplementi di quotidiani: 68% del totale
- Periodici: 32% del totale

La tariffa deriva inoltre dalla somma di due componenti:

- Fissa
- Variabile

Componente fissa

Per determinare la componente fissa della tariffa, viene considerato il numero delle testate aderenti nella categoria di riferimento (Quotidiani/Supplementi; Periodici)

Componente variabile

Per calcolare la componente variabile della tariffa si utilizza una formula che pondera il numero di lettori di ciascuna testata. Per il 2026 si considera il dato di *currency* dell'edizione 2025/I, come rilevato dal *sistema Audipress* e disponibile attraverso le piattaforme di consultazione dei dati.

Ripartizione tra componente fissa e variabile

Quotidiani/Supplementi di quotidiani

- 25% componente fissa
- 75% componente variabile

Periodici

- 40% componente fissa
- 60% componente variabile

Alle testate che non dispongono dei dati di lettura (ad esempio in caso di prima iscrizione) viene attribuito in via provvisoria un numero di lettori calcolato sulla base del numero di lettori di testate già rilevate aventi analoghe caratteristiche. Al momento della disponibilità dei dati effettivi verrà poi effettuato un conguaglio.

L'accesso ai servizi di consultazione del dato è regolato come segue:

- a ciascun Publisher iscritto vengono forniti i dati pubblicati in formato digitale. Di questi dati viene assicurata la diffusione sul mercato per utenti ed agenzie di pubblicità. I dati completi, prodotti dal sistema di rilevazione Audicom Print vengono pubblicati sui canali web istituzionali della società
- i dati pubblicati da Audicom sui propri canali web istituzionali sono liberamente consultabili.
- l'archivio di pianificazione è cedibile da Audicom alle software house accreditate alle condizioni di fornitura ed economiche di cessione che vengono decise, edizione per edizione.

Servizio di distribuzione dati “Audicom Digital & Print”: modalità di calcolo delle tariffe

L’adesione al Servizio di Distribuzione dati “Audicom Digital & Print” è vincolata all’adesione ai servizi “Audicom Digital” e “Audicom Print”.

La tariffa prevista per il servizio «Audicom Digital & Print» considera tre componenti:

- ✓ Valore della tariffa «Audicom Print»
- ✓ Valore della tariffa «Audicom Digital Text» ⁽¹⁾
- ✓ Contributo per la elaborazione del «Audicom Digital & Print»

Tariffa «Audicom Digital & Print» = (Valore della tariffa «Audicom Digital Text» + Valore della tariffa «Audicom Print») * 0,25

È previsto un CAP pari a 30.000 euro per l’attivazione del servizio di distribuzione dati «Audicom Digital&Print», con CAP aggiuntivi applicabili alle prime classi di Active Reach della componente Digital Text secondo criteri di proporzionalità (CAP pari a 5.000,00 euro per AR fino al 5% e pari a 10.000,00 per AR fino al 10%)

NOTE:

⁽¹⁾ nel caso di concessionarie che iscrivono publisher «pure» digital insieme a publisher print il valore della componente «Audicom Digital Text» è calcolato esclusivamente sul valore «digital text» delle testate «print» che fruiscono del dato «Audicom Digital & Print».

-----oooooooooooo-----

UTILIZZATORI PROFESSIONALI – Agenzie Media

Per gli Utilizzatori professionali – Agenzie Media le tariffe sono calcolate sulla base delle dimensioni del *online billing* dell’azienda, ovvero del fatturato pubblicitario online annuale gestito dall’Agenzia per conto dei suoi clienti, dichiarato alla sottoscrizione del contratto e aggiornato annualmente.

Online billing (euro)	Componente fissa (euro)	Componente variabile	Servizio base	Accesso ad Audicom Media View daily, weekly, monthly (euro)
< 1 milione	3.486	(online billing / 1.000.000) * 1.212 euro	Licenza d'uso di Audicom Database Respondent Level Digital (include i dati “Digital&Print” e “Combined&CTV”)	+ 17.435 euro
tra 1 milione e 4 milioni	20.921	(online billing / 1.000.000) * 1.212 euro		Compreso
> 4 milioni	34.867	(online billing / 1.000.000) * 1.212 euro		

Nota: il valore massimo consentito calcolato come somma di [Componente fissa + Componente variabile] è pari a 140.000 euro.

ALTRI UTILIZZATORI (Aziende)

Le tariffe per gli Altri Utilizzatori (Aziende) che desiderano accedere ad Audicom Media View (dati daily, weekly e monthly) e *Audicom Database Respondent Level Digital* seguono le stesse logiche usate per i Publisher iscritti con riferimento all'insieme delle Properties di cui l'Utilizzatore è titolare. La quota minima per accedere ai dati è di 28.500 euro, la quota massima 346.000 euro.

ALTRI UTILIZZATORI (Publisher in CRU)

Questi Soggetti hanno la possibilità di attivare un servizio di accesso diretto per la consultazione dei dati di Audicom Digital.

Servizi inclusi	Modalità di accesso	Tariffa (euro)
- Dati daily, weekly di tutti gli iscritti con SDK Text - dati monthly per tutto il catalogo	piattaforma <i>Audicom Media View</i> (previa attivazione di uno o più account)	18.000
licenza d'uso di <i>Audicom Database Respondent Level Digital</i> (include i dati "Digital&Print" e "Combined&CTV")	<i>analysis & planning tool</i> di software house accreditate (servizio non incluso nell'offerta)	

Note finali su servizi e tariffe

È prevista la possibilità di attuare una revisione annuale delle tariffe per tutte le componenti che determinano il valore della sottoscrizione dei servizi in funzione delle fluttuazioni del mercato pubblicitario e degli investimenti sostenuti dalla medesima per lo sviluppo dei servizi.

Tutti i corrispettivi sono riferiti a 12 mesi solari (gennaio-dicembre) e sono indicati al netto di IVA.

Per l'utilizzo del dato in relazione alla fornitura di alcune tipologie di servizi, potranno applicarsi tariffe specifiche non indicate nei tariffari sopra riportati.

-----oooooooooooo-----

Servizio di distribuzione dati "Audicom ADV Census": modalità di calcolo delle tariffe

Tariffe di accesso al Servizio per gli "Altri Utilizzatori (Aziende)"

Il valore di accesso al Servizio di Distribuzione dati "Audicom ADV Census" è così calcolato:

- Tariffa Fissa di sottoscrizione pari a € 5.000,00 all'anno
- Tariffa Variabile in base alla quantità di settimane attive per campagna/prodotto sulle properties video digitali nel 2025 dichiarata dall'Utilizzatore Professionale (Azienda).

Il valore massimo della tariffa variabile è pari a 200 settimane attive per campagna/prodotto all'anno per un totale di € 50.000,00.

Il valore di ogni singola settimana attiva per campagna/prodotto è pari ad € 250,00, valore che potrà essere ridotto nel secondo semestre 2026 in funzione del numero di aziende che sottoscriveranno il servizio “Audicom ADV Census” nel corso del primo semestre 2026.

Tariffe di accesso al Servizio per gli “Utilizzatori Professionali – Agenzie Media”

Il valore di accesso al Sistema di distribuzione dati “Audicom ADV Census” è pari ad una percentuale pari al 40% del valore di accesso ai dati del Servizio Audicom Digital.

2.3. Metriche

Metriche primarie

“Audicom Digital”, contenuti editoriali:

- Unique Audience
- Page Views
- App Launches
- Legitimate Stream Views
- Time Spent (Text)
- Time Spent (Video)

“Audicom Print”, contenuti editoriali:

- Per i Quotidiani
 - Lettori nel giorno medio
- Per i Periodici
 - Lettori ultimo periodo

“Audicom Digital&print”, contenuti editoriali:

- Utenza unica Digital&Print, derivata dalla somma deduplicata dei lettori “print” e della audience “digital”

“Audicom ADV Census”, contenuti pubblicitari video “digital”

- Stream Views
- Legitimate Stream Views
- Time Spent (Video)

3. RICERCA DI BASE

(in ottemperanza agli adempimenti richiesti dall'Autorità per le Garanzie nelle Comunicazioni nell'ambito delle misure attuative per l'acquisizione, l'elaborazione e la gestione delle informazioni richieste nell'atto di indirizzo sulla rilevazione degli indici di ascolto e di diffusione dei mezzi di comunicazione: delibera 130/06/CSP, art. 6, pubblicato su G.U. 174 del 26/06/2006)

3.1 Dati anagrafici della Società di ricerca

Ipsos S.r.l.

Via Tolmezzo 15, 20132 Milano

Tel + 39 02 36105.1

Fax + 39 02 36105.902

Fax + 39 02 36105.904

Capitale Soc. € 2.000.000

Tribunale 156521 – 3832-21

C.C.I.A.A. 869967

I.V.A. N° 01702460153

Legale Rappresentante Nicola Neri

3.2 Obiettivo della rilevazione

La ricerca ha l'obiettivo di fornire dati relativi alla struttura demosociale delle famiglie residenti in Italia e alle principali dotazioni tecnologiche di interesse per Auditel presenti nelle loro abitazioni principali.

L'indagine di base ha anche ad oggetto le stime sulle dotazioni individuali, ottenute tramite rilevazioni auto-riferite relative al possesso, alla disponibilità (intesa come possibilità di utilizzo, al di là del possesso personale) e all'effettivo utilizzo personale di queste dotazioni (prettamente individuali come ad esempio lo smartphone, o individuali/famigliari come ad esempio il tablet e il Pc).

3.3 Oggetto della ricerca

- Stima ed aggiornamento degli universi di riferimento del panel AUDITEL e di AUDICOM Panel per la rilevazione dei dati delle audience digitali distribuiti da Audicom
- Creazione di un serbatoio di nominativi di famiglie, dal quale attingere per alimentare il panel AUDITEL
- Rilevazione nelle abitazioni principali delle famiglie italiane, della diffusione di dotazioni tecnologiche atte alla fruizione di media e della possibilità di ricezione delle emittenti con diverse modalità/fonti

- Rilevazione auto-riferita del possesso, della disponibilità (intesa come possibilità di utilizzo, al di là del possesso personale) e dell'effettivo utilizzo personale di dotazioni prettamente individuali (come ad esempio lo smartphone) o individual/famigliari (come ad esempio il tablet e il Pc), con l'obiettivo di individuare e quantificare quale parte della popolazione che vive in Italia abbia accesso a internet, con quali modalità specifiche acceda, da quali luoghi, con quali device, con quale frequenza lo utilizzi.

3.4 Metodologia

3.4.1 Consistenza del campione oggetto di indagine

La Ricerca di Base è costituita da una serie continuativa di indagini sull'Universo delle famiglie e degli individui che vivono in Italia (suddivisi in 7 cicli mensili).

Complessivamente la Ricerca di Base è costituita da 20.000 interviste annuali, funzionali a desumere i parametri descrittivi dell'Universo famiglie ed effettuate con campionamento probabilistico, cui si aggiungono le interviste individuali, - funzionali a produrre le stime sulla popolazione che accede ad internet ed effettuate estraendo casualmente un componente a famiglia - ripartite in 10.000 interviste individuali al portavoce familiare (ovvero il componente che risponde alle domande che riguardano la famiglia) e – in base alla resa ottenuta nella Ricerca di base delle edizioni precedenti – massimo 4.000 interviste individuali a un secondo membro della famiglia.

Tali interviste possono essere integrate da campioni di numerosità variabile di casi, finalizzati esclusivamente all'alimentazione del solo Panel Auditel e mirate a specifici segmenti di composizione familiare o ai cittadini stranieri presenti sul territorio italiano.

3.4.2 Modalità di rilevazione ed eventuale margine di errore

La raccolta delle informazioni avviene tramite un unico questionario che contiene le informazioni relative sia alle dotazioni tecnologiche delle famiglie che alle dotazioni individuali.

In particolare, per le dotazioni familiari per quanto attiene la penetrazione delle piattaforme:

- ricevitore digitale terrestre
- ricevitore digitale satellitare
- abbonamenti alla pay-tv

Vengono, inoltre, monitorati il numero di televisori posseduti e le loro caratteristiche, nonché i dati sociodemografici relativi ai singoli componenti la famiglia.

Riguardo alla raccolta delle informazioni sulle dotazioni a livello individuale, il questionario prevede di rilevare possesso, disponibilità (intesa come possibilità di utilizzo, al di là del possesso personale) ed effettivo utilizzo personale di queste dotazioni (prettamente individuali come ad esempio lo smartphone, o individual/famigliari come ad esempio il tablet e il Pc) con l'indicazione di quando è avvenuto l'ultimo accesso alla rete.

Le interviste sono effettuate presso l'abitazione principale delle famiglie dall'intervistato mediante il supporto del personal computer (CAPI – Computer Aided Personal Interviewing).

Fanno eccezione le interviste individuali ai secondi membri estratti limitatamente ai casi in cui la persona da intervistare non sia in casa o non sia al momento disponibile per l'intervista: in questo caso le interviste vengono infatti svolte in un secondo momento al telefono (dal centro telefonico Ipsos in modalità CATI - Computer Aided Telephone Interviewing o dall'intervistatore che ha svolto l'intervista principale in modalità CAPI - Computer Assisted Personal Interviews).

In entrambi i casi il supporto del personal computer consente un controllo ed un cleaning in tempo reale delle informazioni raccolte.

Prima di effettuare le interviste tutti gli intervistatori coinvolti nella rilevazione sono stati addestrati tramite apposite riunioni di briefing e dotati di materiale di supporto e manuali specifici per il corretto svolgimento dell'indagine.

Durante lo svolgimento della rilevazione, tutte le fasi vengono costantemente monitorate al fine di garantire la corretta esecuzione operativa della ricerca.

Una volta completata la raccolta sul campo, le informazioni vengono trasmesse all'istituto per via telematica e subiscono le operazioni di contabilizzazione e controllo di qualità tramite verifica telefonica con l'intervistato.

Il file di lavoro viene settimanalmente sottoposto ad un ulteriore processo di cleaning, che si basa sull'immissione manuale di controlli di coerenza o di conformità agli standard richiesti.

Il margine di errore relativo ai risultati della ricerca (livello di significatività del 95%) è compreso fra +/- 0,14 % e +/- 0,69 per i valori percentuali relativi al totale degli intervistati (20.000 casi).

Adeguamento metodologico della Ricerca di Base all'emergenza sanitaria causata dal Covid-19

In considerazione del possibile protrarsi di effetti negativi connessi al COVID-19, l'indagine di base è così strutturata:

- Mantenimento frame indirizzi e primo contatto effettuato sempre da parte della rete di intervistatori face to face → Dotando la rete degli opportuni dispositivi di sicurezza è possibile mantenere questi due capisaldi metodologici dell'indagine, a garanzia dell'effettiva probabilisticità del campione
- Tecnica di intervista multimode → È prevista la possibilità di ricorrere alla consueta intervista domiciliare per chi accetta, mentre per chi resiste a far entrare in casa l'intervistatore è prevista la possibilità di effettuare l'intervista all'esterno dell'abitazione (tipicamente sul pianerottolo), chiedendo di parlare con responsabile tecnologico o procedendo con un qualsiasi componente della famiglia a cui viene lasciata la scheda dotazioni con la richiesta che venga compilata dal responsabile tecnologico (immediatamente se è in casa o in un secondo momento).

3.4.3 Universo di riferimento

La Ricerca di Base fornisce due tipi di statistiche: statistiche relative alle famiglie (numero di televisori posseduti, attrezzature e dotazioni tecniche) e statistiche relative agli individui (informazioni socio-demografiche rilevate per ogni componente della famiglia, oltre a possesso, disponibilità ed utilizzo personale di dotazioni prettamente individuali o individual/famigliari, con l'indicazione di quando è avvenuto l'ultimo accesso alla rete, tramite rilevazione auto-riferita presso un componente a famiglia estratto casualmente).

Gli universi di riferimento esaminati (oggetto della ricerca) sono perciò differenziati e cioè:

A) Universo di riferimento famiglie:

Famiglie che vivono sul territorio italiano, comprese le famiglie interamente composte da stranieri. Per famiglia si intende l'insieme di persone che vivono nella medesima abitazione, indipendentemente da vincoli di parentela/affettività o mutuo sostegno economico. Fonte per l'universo famiglie:

- per il campionamento → Istat – Bilancio Demografico Istat al 31 dicembre 2022 (26.400.326), che però si basa su una definizione di famiglia che esclude la semplice coabitazione

Poiché la definizione di famiglia adottata dall'indagine differisce dalla definizione di famiglia Istat, la Ricerca di Base mutua quanto fatto da Auditel a partire dalla fine del 2014 (quando ancora Auditel e Audiweb Srl – prima della fusione per unione con Audipress Srl in Audicom – svolgevano indagini di base separate e distinte), procedendo, in fase di analisi, ad una stima autonoma delle famiglie sulla base dei dati campionari.

- per la ponderazione → Dati ricavati ricorrendo al metodo di 'traduzione' degli universi, che prevede il confronto tra la composizione familiare da anagrafe e la composizione familiare da dichiarato delle famiglie della Ricerca di base Auditel 2015. In particolare, il confronto ha riguardato le interviste realizzate con il metodo indirizzi nel 2015 (integrate dalle interviste su quota per il campione stranieri realizzate sempre nel 2015, al fine di avere a disposizione una base casi più robusta), per le quali sono state recuperate le informazioni anagrafiche.

I dati di confronto tra situazione familiare dichiarata e situazione familiare registrata in anagrafe sono stati utilizzati per ricondurre i dati pubblicati Istat (che escludono le coabitazioni non connotate da affettività/parentela o mutuo sostegno) a dati omogenei alla definizione di famiglia adottata da Auditel e mutuata nella Ricerca di Base.

La base casi in analisi è complessivamente costituita da 19.098 interviste.

B) Universo di riferimento individui:

Individui in famiglie che vivono sul territorio italiano.

Fonte per l'universo individui: Istat – Bilancio Demografico popolazione residente al 31 dicembre 2023 (www.demo.istat.it) (58.971.230)

3.4.4 Metodo di campionamento

Il disegno campionario utilizzato nella Ricerca di Base, che corrisponde a quello adottato dalla Ricerca di base Auditel dal 2015, si basa sull'utilizzo di uno schema di campionamento a quattro stadi, di cui il primo stratificato.

- L'unità al primo stadio di campionamento è il Comune (selezione PPS)
- L'unità al secondo stadio è la sezione elettorale (per i comuni di oltre 10.000 abitanti) o l'aggregazione di sezioni censuarie (per i comuni fino a 10.000 abitanti), che fungono unicamente da agglutinatori territoriali (selezione SRS)
- L'unità al terzo stadio è il civico (selezione PPS)
- L'unità finale è la famiglia domiciliata nell'unità abitativa, scelta tramite la selezione casuale dalla lista delle unità abitative all'interno della sezione elettorale estratta/dell'aggregato di sezioni censuarie estratto (selezione SRS)

La stratificazione delle unità primarie prevede:

- il campionamento certo di tutti i capoluoghi di provincia e dei comuni superiori agli 80 mila abitanti (unità autorappresentate) che costituiscono uno strato a sé stante ed assorbono un numero di interviste proporzionale alla loro dimensione (per dimensione si intende il numero di famiglie che risiedono all'interno di quel comune o di quello strato)
- la stratificazione dei restanti comuni per provincia e per ampiezza centro sulla base del numero di residenti (l'allocazione delle interviste, tuttavia, è basata sul numero di famiglie). Entro ogni strato i comuni sono estratti tramite metodo PPS (Probability Proportional to Size), in base al numero di famiglie residenti nel singolo comune.

Le informazioni relative all'indirizzo (civico e unità abitativa) vengono estratte dalla banca dati catastale dell'Agenzia del Territorio.

Per i comuni non inclusi nel DB catastale dell'Agenzia del Territorio il campionamento viene effettuato tramite i viari cittadini.

Estensione territoriale: il campione finale è di 1.323 comuni, compresi i capoluoghi di provincia.

3.4.5 Ponderazione ed espansione dei risultati

Grazie all'adozione di universi di riferimento che garantiscono una perfetta corrispondenza tra famiglie e individui, è stato possibile ponderare i risultati della Ricerca di Base mediante calibrazione, il processo di ponderazione che consente di governare contemporaneamente i parametri individuali e quelli famigliari, facendo sì che la media dei pesi individuali dei componenti di una stessa famiglia sia pari al peso della famiglia (in questo caso la dimensione campionaria ha consentito di aggiungere il vincolo che il peso di ciascun componente sia esattamente pari al peso della famiglia di appartenenza).

La calibrazione garantisce di ottenere stime perfettamente coerenti tra famiglie ed individui a qualsiasi livello di dettaglio si voglia giungere.

Imposizione dei parametri famigliari

Si procede ad imporre i seguenti parametri universo, ricavati tramite il metodo di ‘traduzione’ degli universi, che prevede il confronto tra la composizione famigliare da anagrafe e la composizione famigliare da dichiarato delle famiglie della Ricerca di base Auditel 2015:

- la distribuzione delle famiglie per regione incrociata per ampiezza centro
- la distribuzione delle famiglie per numero di componenti, incrociata per regione e incrociata per ampiezza centro
- la distribuzione delle famiglie tra famiglie di soli italiani, famiglie miste e famiglie di soli stranieri
- la distribuzione per numero di componenti, separatamente per le famiglie di soli italiani e le famiglie con almeno uno straniero

Nel caso della ponderazione su base trimestrale, al file famiglie vengono direttamente imposti per ponderazione anche alcuni dati puntuali, certificati, forniti dagli Editori e in particolare:

- numero totale di abbonati alla pay-tv Sky
- distribuzione per area geografica degli abbonati alla pay-tv Sky

Imposizione dei parametri individuali

Si procede ad imporre i seguenti parametri, ricavati dal Bilancio demografico Istat al 31 dicembre 2023 ad eccezione della ripartizione degli individui per tipologia di famiglia (ricavata con il metodo di ‘traduzione’ degli universi analogamente ai parametri famigliari):

- la distribuzione degli individui per parametri territoriali (regione per ampiezza centro, Provincia)
- la distribuzione degli individui per sesso per classi di età
- la distribuzione degli individui per età e ampiezza centro
- limitatamente agli individui stranieri, distribuzioni a marginale per sesso, età, aggregazioni di nazionalità, area geografica e ampiezza centro
- la distribuzione degli individui per tipologia di famiglia (di soli italiani, di soli stranieri, italiani in famiglie miste, stranieri in famiglie miste)

Nel caso della ponderazione su base trimestrale, ai record individuali viene inoltre imposto il numero di individui 4+ in famiglie abbonate a Sky, ottenuto applicando al dato Sky - imposto in ponderazione a livello famiglie - il numero medio di componenti delle famiglie Sky risultante dalla Ricerca di Base cumulando ogni volta le ultime 5 wave elaborabili.

3.5 Periodo di rilevazione

L'indagine è costituita da 7 rilevazioni mensili. Di seguito il calendario di rilevazione per singola wave.

1° wave F2F 16-1-2025 / 1-3-2025

2° wave F2F 17-2-2025 / 15-4-2025

3° wave F2F 1-4-2025 / 30-5-2025

4° wave F2F 15-5-2025 / 13-7-2025

5° wave F2F 14 / 31-7-2025 e 28-8-2025 / 7-10-2025

6° wave F2F 24-9-2025 / 20-11-2025

7° wave F2F 3-11-2025 / 23-12-2025 e 27 / 30-12-2025

Mentre gli universi vengono aggiornati una volta all'anno per quanto riguarda le principali caratteristiche socio-demografiche della popolazione e una volta ogni tre mesi per quanto riguarda le condizioni di ricezione delle diverse piattaforme di trasmissione e l'accesso ad Internet, le interviste della Ricerca di Base hanno effetto immediato sulla sola rilevazione degli ascolti TV in quanto i nominativi delle famiglie intervistate entrano a far parte del Data Base delle famiglie da contattare per il reclutamento nel campione meterizzato Auditel.

Per il prossimo aggiornamento degli universi delle principali caratteristiche sociodemografiche, previsto per il prossimo agosto, si utilizzeranno le nuove informazioni della Ricerca di Base unitamente ai dati ISTAT (se non dovessero essere pubblicati dei nuovi aggiornamenti, rimarranno quelli adottati attualmente come universo di riferimento).

3.6 Indirizzo web di pubblicazione della metodologia

La metodologia utilizzata per l'esecuzione della Ricerca di Base è disponibile presso la sede di Auditel (Via Larga 11, Milano) o la sede operativa di IPSOS (Via Tolmezzo 15, Milano).

4. AUDICOM DIGITAL PANEL

4.1. Dati anagrafici delle Società di ricerca

The Nielsen Media Italy S.r.l.

Centro Direzionale Milanofiori, Strada 6, Palazzo A11/12/13, Assago (MI)

Legale Rappresentante Luca Bordin

Telefono: 02-32118001

indirizzo PEC: nielsenmediaitalysrl@legalmail.it Cap.

Soc. € 8.184.711,00

Registro Imprese Milano Monza Brianza Lodi n° di iscrizione 11227560965 Trib.

Milano/R.E.A. Milano n.2588448

Cod.Fisc. e P.IVA 11227560965

4.2. Obiettivo della rilevazione

La Ricerca di Base, utilizzando la metodologia del questionario quantitativo, può analizzare fenomeni “macro” per i quali l’intervistato possa dare un’affidabile risposta. La rilevazione dei dettagli dell’utilizzo di Internet (siti e sezioni di siti visitati, tempo speso, etc.) non può perciò essere rilevata tramite dichiarazione dell’intervistato.

Per ottenere tali informazioni occorre disporre di un campione continuativo, panel, di individui sui quali sia possibile effettuare una rilevazione tecnica dell’effettivo comportamento di navigazione.

La componente Panel di Audicom Digital si avvale di appositi software meter installabili su PC/Mac, tablet e smartphone che consentono un monitoraggio continuativo dell’attività online dei panelisti (siti web e App mobile su sistemi operativi iOS e Android).

4.3. Metodologia

4.3.1. Software Meter per la rilevazione PC (Desktop e Laptop)

La soluzione tecnica per misurare l’audience dei computer è il meter NetSight gestito da Nielsen. Si tratta di un software meter (compatibile con PC e Mac) disegnato per rilevare il comportamento online degli utenti che navigano da computer. Per garantire una maggiore accuratezza dei dati, la rilevazione viene condotta distinguendo gli utenti che utilizzano lo stesso computer in famiglia. La procedura prevede che il panelista possa effettuare il login all’inizio di ogni sessione o dopo un periodo di inattività, inoltre vengono impiegati sistemi di modelling.

I dati ricevuti dal meter NetSight gestito da Nielsen vengono elaborati per garantire che vengano conteggiate solo le attività di navigazione rilevanti per la misurazione delle attività online dei panelisti. Il trattamento si basa su una serie di regole di accredito (crediting) sviluppate in collaborazione con il JIC che aiutano a fornire metriche rappresentative del reale comportamento degli utenti che utilizzano il computer. Il meter viene aggiornato regolarmente per mantenere la compatibilità con tutti i principali browser e con l'evoluzione dei rich media. Gli aggiornamenti del meter vengono inviati automaticamente a tutti gli individui inclusi nel campione.

Il meter NetSight ha la capacità di misurare l'attività all'interno di una pagina web che risulti essere 'in focus' in un dato momento - ovvero visualizzato nel tab attivo del browser. Questo tipo di misurazione garantisce una rilevazione accurata della durata e delle metriche del tempo speso sulla pagina e consente di superare i problemi di errata attribuzione del traffico dovuti alle finestre ridotte a icona e alla navigazione del browser a schede.

Le regole di crediting delle pagine viste e del tempo speso possono essere soggette ad aggiornamenti costanti che riflettono la natura dinamica delle tecnologie web; il principio dominante resta sempre quello di dare una rappresentazione accurata della navigazione dell'utente, attribuendo correttamente audience, durata e volumi di consumo (Pagine Viste) ai relativi siti web che li hanno generati. La soluzione qui descritta include le regole più aggiornate disponibili al momento.

Le regole di crediting di Pagine Viste e Durata si basano sulla misurazione degli eventi Web Browser Standard: questo approccio permette di identificare gli eventi validi e di scartare il traffico non valido, permettendo un elevato grado di precisione nella rilevazione delle Pagine Viste e della Durata.

4.3.2. Raccolta dei dati panel

Il meter NetSight è in grado di misurare il traffico di rete, i dati immessi dall'utente, la navigazione web, il login e logout degli utenti e altri aspetti dell'utilizzo del computer, come l'output audio, quest'ultimo tuttavia, non è utilizzato nella metodologia qui descritta.

Tutti i dati sottoposti a misurazione sono raccolti nei log file. I log file vengono archiviati e codificati sulla macchina del panelista per consentire una rilevazione continua e vengono raccolti e conservati anche quando il computer è offline. Una volta online, i log file vengono inviati all'infrastruttura di raccolta Nielsen e infine inoltrati per l'elaborazione.

L'invio del log file non ha nessun vincolo di dimensioni, e avviene in base ad uno specifico intervallo di tempo. Tale intervallo è specifico per ogni singola configurazione e solitamente è intorno ai 10 minuti. Il log file viene inviato allo scadere dell'intervallo temporale stabilito anche se il limite di 512 KB non è stato raggiunto. In questi casi, il file in uso viene inviato e si procede alla creazione del successivo.

Il meter NetSight registra le seguenti informazioni:

CATEGORIA	INFORMAZIONI	NOTE
Identificazione Utente	ID Utente tramite finestra di accesso del meter	Specificamente nel caso di nuclei familiari con più partecipanti al panel. Il prompt viene presentato all'utente all'inizio di ogni sessione. Una sessione inizia all'avvio della macchina, al momento dell'accesso o alla prima attività dell'utente dopo 30 minuti di inattività.
	Nome utente di Windows	Viene catturato solo al momento dell'inserimento durante la fase di login iniziale. Viene acquisito solo il nome utente e non la password.
	Password di riconoscimento Biometrico	Nelle famiglie con più partecipanti al panel, viene monitorata la velocità di scrittura delle parole di uso comune come "il", ".com", "e", per supportare la fase di identificazione dell'utente.
	Nome sul riconoscimento biometrico.	Nelle famiglie con più panelisti vengono tracciati i primi quattro caratteri del nome utente più il cognome per identificare l'utente di ogni sessione.
	Nome utente nella URL	Se il nome utente è presente nell'URL (es. www.hotmail.com/userid/susan/jds?), viene ricercato in fase di post-elaborazione dei dati.
Informazioni sulla Sessione	Eventi di attività del computer	Registrazione delle volte in cui il panelista è attivo.
	Attività dell'utente su tastiera, mouse e altri input	Viene registrato il Timestamp per ogni transizione di un dispositivo dalla modalità "Attivo" a "Inattivo" o viceversa. Se non viene effettuato alcun accesso tramite il dispositivo (es. tastiera) per 60 secondi, il dispositivo verrà considerato Inattivo. 30 minuti di inattività o stand-by terminano la sessione. Questo periodo di inattività è configurabile.
Utilizzo del web	URL di siti Web pubblici e sicuri e ora di accesso	Gli URL (Uniform Resource Locator), ad es. www.google.com digitati dall'utente vengono catturati, così come gli URL di referrer e di reindirizzamento.
	Informazione sulla finestra del browser	Per ogni pagina visualizzata vengono acquisite informazioni relative alla dimensione e alla posizione del browser.
	Contenuto nuova pagina v. contenuto corrente	URL aperto come nuova pagina o come pagina corrente, viene raccolto. (1Q'09)
	Cookie	Il meter implementa l'acquisizione dei cookie in questo modo: viene catturato solo un numero limitato di cookie per URL. (con limite predefinito di 10 cookie). La priorità viene data ai cookie configurati per essere monitorati.
	TAG per la misurazione censuaria	I tag utilizzati per la misurazione censuaria (es.: SDK) vengono acquisiti come parte del traffico di navigazione.

4.3.3. Software Meter per la rilevazione Mobile (Smartphone e tablet)

La soluzione per la misurazione passiva delle audience mobile avviene attraverso un'app (meter), scaricabile sui dispositivi mobili, che raccoglie dati relativi all'utilizzo di siti web e app, su piattaforme smartphone e tablet Android e iOS che utilizzano connettività cellulare e Wi-Fi.

La soluzione di misurazione per iOS prevede che l'applicazione (meter) configuri una connessione di rete privata virtuale on device in grado di rilevare il traffico di rete del dispositivo.

Utilizzando le informazioni contenute nello userAgent (per le versioni precedenti a iOS 15) e i pattern delle URL (per tutte le altre versioni), il meter è in grado di determinare tramite quale app sta avvenendo l'attività online.

Nel caso della misurazione dei dispositivi Android, il meter è costituito da un'app scaricabile dal Google Play Store che, dopo la configurazione, rimane attiva in background sul dispositivo del panelista. Per la rilevazione web browsing su Android, il meter configura una connessione di rete privata virtuale on device in grado di annotare il traffico di rete delle app dei principali browser. Per la rilevazione dell'utilizzo di tutte le altre app, il meter si affida ad appositi strumenti del sistema operativo Android che rendono possibile la rilevazione delle attività delle app in foreground

I panelisti vengono forniti di istruzioni dettagliate su come scaricare e configurare il meter sul proprio dispositivo e sono monitorati in modo continuativo per accertarne l'attività. Questo aspetto è particolarmente importante per la misurazione mobile poiché il ciclo di vita dei dispositivi mobili è molto più breve del ciclo di vita dei computer e gli aggiornamenti sono relativamente frequenti.

Tutti i dati vengono raccolti in tempo reale a livello di URL e app e inviati ai server di Nielsen che procede alla rimozione del traffico non valido e "irrilevante" (es. URL tecniche che non rientrano nel perimetro di misurazione) e alla successiva formattazione dei dati. Alla stregua del meter NetSight, il meter mobile viene regolarmente aggiornato per garantire la compatibilità con gli aggiornamenti dei sistemi operativi. Gli aggiornamenti sono inviati in modo automatico ai device mobile sottoposti a metering.

4.4. Campione oggetto di indagine

4.4.1. Single digital recruitment

I panelisti, quale che sia il dispositivo da cui si collegano, possono unirsi al campione attraverso un sito di reclutamento unico. Con un unico set di credenziali, i panelisti possono gestire i dati demografici del loro nucleo familiare e i device supportati dalla misurazione.

In generale, il sito adotta un concetto di "member area", caratterizzato da dashboard moderne e di facile utilizzo.

Il sito consente il reclutamento di un campione con le seguenti caratteristiche:

- permette ai panelisti di aderire ai singoli panel desktop, smartphone, tablet oppure ad una combinazione di essi, entro i limiti delle necessità di gestione del campione;
- permette ai panelisti di unirsi a tutti i campioni con più device dello stesso tipo;

- permette la rilevazione delle persone di età 4+ nel panel desktop;
- permette la rilevazione nel panel mobile 13+.

4.4.2. Reclutamento dei Minori per la rilevazione mobile

Il reclutamento e la misurazione del comportamento online dei minori costituisce un problema complesso, soprattutto a causa del necessario livello di protezione che deve essere integrato nel sistema per un campione così delicato.

Le principali aree di attenzione quando si parla di minori riguardano la tutela della loro privacy e dei dati raccolti e, soprattutto, la corretta acquisizione del consenso.

Ai sensi di legge, i minori non hanno la capacità giuridica di sottoscrivere personalmente un contratto, e di assumerne i relativi diritti e doveri. Ciò pone criticità non solo nel processo di registrazione, ma anche nei normali processi di gestione del panel: anche se il consenso è stato inizialmente acquisito, qualsiasi altra azione, incluso il riscatto dei punti, presenta sfide da risolvere.

Nielsen recluterà minori per tutti i campioni attraverso la mediazione di un adulto che, per i termini del contratto, deve avere il diritto di registrare i dispositivi utilizzati da altri e fungere da tutor per qualsiasi minore che lo utilizzi. Il processo è disegnato per richiedere all'utente di eseguire azioni che confermano il consenso durante il percorso.

Ciò consente di reclutare direttamente nuove famiglie con minori e anche di aggiungerli alla misurazione mobile per famiglie che eventualmente non li avessero inizialmente inclusi nella misurazione.

Dimensione del Panel “meterizzato”

- Panel TOTAL: 26.000 panelisti eleggibili ed attivi
- Panel Probabilistico: 3.900 panelisti eleggibili ed attivi
- Panel Mobile: 10.400 panelisti eleggibili ed attivi (di cui 1.400 con tablet)
- Panel Single Source: all'interno del campione Computer sono inclusi 2.000 panelisti “single-source” (ovvero che fanno parte anche del Panel Mobile).

4.5. Ponderazione ed eleggibilità per i dati panel PC e Mobile

4.5.1. Processo di ponderazione per la componente Computer

La produzione dei dati si basa su diversi campioni (panel computer e panel mobile) e metodologie di campionamento, una probabilistica e l'altra Online (Convenience).

Il processo di ponderazione per la componente Computer si compone di due fasi:

- ponderazione del campione probabilistico rispetto agli universi della Ricerca di Base
- ponderazione del campione totale (probabilistico + online) rispetto alla somma dei pesi, utilizzando target comportamentali per correggere eventuali distorsioni del campione reclutato online.

Il processo verifica, alla fine della convergenza, se la somma dei pesi corrisponde alle stime degli universi da Ricerca di Base per categoria. La somma dei pesi deve essere entro l'1% della Ricerca di Base. È ammissibile un massimo di 100 iterazioni.

Ai dati viene applicato un trimming che permette di evitare di computare pesi troppo alti e/o troppo bassi: nel processo iterativo di rim-weighting, se un peso ha un valore inferiore a 1/4 o superiore a 4 rispetto al peso medio, tale peso viene corretto al valore di soglia. Ciò migliora la convergenza dei pesi nelle varie iterazioni evitando pesi estremi.

Fase 1: Ponderazione campione probabilistico

Il campione probabilistico è ponderato in base a caratteristiche demografiche specifiche per un determinato mercato. Nel caso dell'Italia vengono riportate di seguito tali caratteristiche utilizzate per la ponderazione e il numero minimo di gruppi di controllo per ciascuna di esse. I gruppi demografici oggetto di controllo sono solitamente stabili nel tempo, tuttavia la definizione di tali gruppi potrà essere soggetta a cambiamenti resi necessari da aggiornamenti della metodologia di ponderazione o da specifiche richieste del JIC.

Caratteristica demografica	Numero minimo di gruppi di controllo
Età per genere	7
Istruzione	3
Condizione lavorativa	6
Dimensioni nucleo familiare	4

Fase 2: Ponderazione totale del campione

Il campione totale è ponderato in base a target demografici e comportamentali:

- Target demografici: sono simili alla fase 1 del processo di ponderazione ma non necessariamente identici. Anche in questo caso i gruppi demografici oggetto di controllo sono solitamente stabili nel tempo, tuttavia la definizione di tali gruppi potrà essere soggetta a cambiamenti resi necessari da aggiornamenti della metodologia di ponderazione o da specifiche richieste del JIC

Caratteristica demografica	Numero minimo di gruppi di controllo
Età per genere	7
Istruzione	3
Condizione lavorativa	6
Dimensioni nucleo familiare	4
Lifestage	4
Regione	5

- Target comportamentali: è un tipo di ponderazione particolarmente utile per correggere le distorsioni dei panel reclutati tramite siti online che usano traffico incentivato/ paid-to-click. Aggiungendo questi siti come target comportamentali (in pratica, ponderando l'intero Panel sulla reach% di tali siti nel Panel probabilistico) tale impatto viene ridotto.

- o Controlli comportamentali aggregati:

- Decili basati sul Totale delle Pagine Viste
- Decili basati sul Tempo Speso Totale
- Decili basati sul Totale delle Sessioni

Ognuno degli elementi succitati ha 11 variabili di ponderazione (1° decile, 2° decile, etc + Inattivo)

- o Controlli comportamentali a livello di brand (laddove siano disponibili le stime degli universi) o delle seguenti entità (basati sulla Unique Audience:

- Entertainment- Gambling&Sweepstakes (Categoria)
- Entertainment - Video & PC Gaming/Online Gaming (Categoria)
- Finance/Insurance/Investment - Credit Cards (Categoria)
- Finance/Insurance/Investment-Loans (Categoria)
- Commerce - Coupons & Rewards (Categoria)
- Commerce - Free Merchandise (Categoria)
- Search Engines/Portals & Communities - Member Portals & Communities (Categoria)

Ognuno dei gruppi succitati ha 3 tipi di variabili di ponderazione (Visitato, Non Visitato e Inattivo)

I gruppi comportamentali possono essere oggetto di revisione a seguito di aggiornamenti della ponderazione o per cambiamenti avvenuti nella definizione delle entità, dei brand o dei siti più visitati.

Per evitare la non convergenza dei processi di ponderazione, vengono applicate delle collapsing rule. Se prendiamo ad esempio la categoria "Entertainment- Gambling & Sweepstakes", le possibili variabili di ponderazione sono "Visitato", "Non visitato" e "Inattivo". Se il controllo della proporzione tra Panel e Universo fallisce, la variabile di ponderazione viene collassata perché, se i panelisti vengono combinati nelle categorie "visitato" e "non visitato", il controllo di ponderazione risulterebbe inefficace. Queste collapsing rule vengono attivate per una determinata variabile nel controllo di pesatura se non si verifica almeno una delle due seguenti condizioni:

1. Il rapporto Panel/Universo è al di fuori dell'intervallo compreso tra 0,25 e 4. Il rapporto Panel/Universo viene calcolato come il rapporto tra la percentuale di individui in quel gruppo di ponderazione nel panel e la percentuale di individui in quel gruppo di ponderazione nella stima dell'universo. Se il rapporto tra Panel/Universo è superiore a 4 o inferiore a 0,25, il gruppo di ponderazione viene collassato secondo le collapsing rules.
2. Se la dimensione del campione per un particolare controllo di ponderazione è inferiore a 16, soglia ritenuta ottimale da Nielsen per la specifica dimensione del panel, i controlli di ponderazione vengono collassati in base alle collapsing rule.

4.5.2. Processo di ponderazione per la componente Mobile

La metodologia mobile utilizza una procedura basata sul RIM weighting che garantisce, per ogni controllo di pesatura, che la somma dei pesi rientri nell'1% di quanto riportato dalla stima dell'universo per le stesse caratteristiche. Anche in questo caso i dati sono sottoposti a trimming per evitare che i pesi calcolati siano eccessivamente alti o eccessivamente bassi.

Nel processo iterativo del processo di RIM weighting, se un peso ha un valore inferiore a $\frac{1}{4}$ o superiore a 4 volte il peso iniziale, viene corretto fino a un valore di soglia.

I controlli di ponderazione per la componente mobile del panel sono elencati di seguito:

Caratteristica demografica	Numero minimo di gruppi di controllo
Smartphone per Età per Genere	8
Sistema Operativo per Genere	4
Sistema Operativo per Età	10
Tablet per Età per Genere	8
Device Type	3

Tali controlli di ponderazione sono soggetti a modifiche se si verificano evoluzioni del mercato e/o qualora lo richieda un aggiornamento del sistema di ponderazione.

4.5.3. Criteri di eleggibilità

- **Criteri di eleggibilità per il panel desktop**

I criteri di eleggibilità prevedono criteri stringenti relativi all'attività del campione e alle caratteristiche dei panelisti. Ciò consente una rappresentazione ottimale dell'attività degli utenti attivi. I criteri di eleggibilità considerano due fattori chiave:

- Criteri di qualificazione
 - presenza di almeno un adulto maggiorenne che possa essere designato come "capofamiglia"

- dimensione del nucleo familiare inferiore a 15 individui
- età e genere validi
- Criteri fattuali (ad esempio: il meter è stato installato da abbastanza tempo prima del periodo di rilevazione? Quando è stato inviato l'ultimo log da questo meter? Ecc.)

I criteri di eleggibilità consentono un attento controllo dell'eleggibilità a livello individuale, aggiungendo condizioni specifiche.

• Criteri di eleggibilità per il panel mobile

I partecipanti al panel mobile sono considerati attivi se hanno inviato un log negli ultimi 30 giorni (incluso, nel caso di dispositivi Android, l'attività offline).

I criteri di eleggibilità per i sistemi operativi Android e iOS considerano due fattori principali:

- Criteri di qualificazione
 - età 13-74 anni (*la quota 13-17 anni verrà gradualmente implementata nel corso del 2026*)
 - età e genere validi
- Criteri fattuali (ad esempio: il meter è stato installato da abbastanza tempo prima del periodo di rilevazione? Quando è stato inviato l'ultimo log da questo meter? Ecc.)

Si specifica che il limite di età a 74 anni è dettato dalla mancanza, allo stato attuale, di informazioni dalla ricerca di base (RdB) circa l'utilizzo del mobile da parte della popolazione al di sopra di questa fascia di età. Qualora la RdB fosse aggiornata per colmare questa lacuna, Nielsen provvederà ad aggiornare di conseguenza i criteri di eleggibilità per il panel mobile, a parità di dimensione del campione in oggetto.

4.5.4. Controlli antifrode sui panelisti

Gli sviluppi tecnologici rendono più facili i tentativi di creazione di profili falsi il cui unico scopo è quello di ottenere incentivi in modo fraudolento.

Diversi team sono coinvolti nell'identificazione dei panelisti fraudolenti e i controlli sono continuamente aggiornati. Il processo include controlli preventivi, che precedono la registrazione al panel come l'utilizzo di "fraud scoring" forniti da terze parti, verifiche via SMS e verifiche dell'indirizzo. Dopo che la registrazione è avvenuta vengono inoltre applicate una serie di verifiche automatiche di identificazione proprietarie e vari controlli manuali dei dati panel.

In aggiunta ai controlli standard Nielsen lavora con un partner specializzato, Ehawk, nell'identificazione di comportamenti sospetti e indicativi di bot, spam e utenti fraudolenti, al fine di bloccare questo tipo di contatti già in fase di registrazione.

È estremamente importante che la denominazione esatta dei partner, in special modo di Ehawk, e delle liste utilizzate rimangano confidenziali per proteggere l'efficacia di queste misure e per questo motivo i dettagli non dovranno essere resi pubblici per nessun motivo,

salvo i casi previsti dall'Accordo e/o laddove necessario per conformità alla Legge Applicabile e fermo comunque restando che Audicom dovrà avere contezza dei partner e delle liste utilizzati da Nielsen.

4.6. Utilizzo del campione convenience digitale nella ricerca stampa CAWI

Verrà approntato un sistema per inviare ai panelisti digitali convenience attivi un invito a partecipare alla ricerca relativa alla stampa. In generale, il processo ha inizio con la selezione di un terzo del campione digitale convenience attivo alla vigilia di una wave CAWI. La selezione sarà congegnata in modo da impedire una doppia selezione nell'anno. Ai panelisti selezionati per ogni specifica wave, Nielsen invierà una email che invita alla partecipazione alla ricerca print, in cambio di specifici incentivi. A quel punto, il panelista dovrà semplicemente selezionare un collegamento per poter raggiungere la piattaforma che eroga il sondaggio.

4.7. Periodo di rilevazione

Rilevazione continuativa

4.8. Indirizzo web di pubblicazione della metodologia

www.audicom.net

5. AUDICOM CENSUS

5.1. Dati anagrafici della Società di ricerca

The Nielsen Media Italy S.r.l.

Centro Direzionale Milanofiori, Strada 6, Palazzo A11/12/13, Assago (MI)

Legale Rappresentante Luca Bordin

Telefono: 02-32118001

indirizzo PEC: nielsenmediaitalysrl@legalmail.it

Cap. Soc. € 8.184.711,00

Registro Imprese Milano Monza Brianza Lodi n° di iscrizione 11227560965

Trib. Milano/R.E.A. Milano n.2588448

Cod.Fisc. e P.IVA 11227560965

AUDITEL S.r.l.

Sede legale: Via Larga 11 - 20122 Milano

P.I./C.F. e Iscrizione R.I. Milano: 07483650151

Numero REA Milano 1164218

Capitale Sociale: euro 300.000 i.v.

Legale Rappresentante: Lorenzo Sassoli de Bianchi

5.2. Obiettivo della rilevazione

L'implementazione su base volontaria di SDK – che si attivano fruendo delle pagine web, delle mobile App e dei player video che li contengono e che inviano informazioni ai server di raccolta dei dati – consente di:

- rilevare in modo oggettivo e completo (censuario) i dati volumetrici (pagine viste, stream views, tempo speso, ...)
- contribuire – grazie al sistema DCR ID – alla generazione di dati volumetrici e di audience, profilati per genere ed età, riferiti al singolo giorno / alla singola settimana (servizio: Audicom Daily/Weekly) che – insieme ai dati Audicom Panel – informano il processo di produzione di Audicom Database RL Digital.

5.2.1. SDK Text – caratteristiche e funzionamento

(a cura di The Nielsen Media Italy s.r.l.)

Il software development kit (SDK) consente di integrare gli strumenti di misurazione nei contenuti digitali statici/testuali indipendentemente dal device utilizzato per la fruizione del contenuto.

L'SDK è disponibile per:

- Browsers (Javascript)
- Applicazioni iOS (libreria nativa)
- Applicazioni Android (libreria nativa)

La seguente tabella descrive nel dettaglio le modalità di rilevazione eseguite dall'SDK Text:

	METRICHE	CONDIZIONI	SPECIFICHE
		Load Condition	La pagina deve essere completamente ricaricata, con browser / tab in focus. In caso di inserimento diretto del Tag/SDK (ossia senza utilizzo di tag manager) il beacon è posizionato in fondo alla pagina
		Focus	Browser e tab devono essere in focus
		Autorefresh	Ogni refresh della pagina effettivamente visualizzata (pagina completamente caricata e browser/tab in focus = una nuova Pagina Vista) Gli autorefresh che avvengono out of focus non sono conteggiati. Il conteggio inizia quanto browser/tab sono in focus. Ogni pagina web visualizzata da Computer/Smartphone/Tablet via web browsing ma richiamata da una funzione di autorefresh con timeout inferiore ai 300 secondi non deve presentare il TAG /SDK ossia non deve generare una chiamata ai server Nielsen per il conteggio della pagina.
	Page Views		

Browser			La compliance a questa regola non è generata in automatico da SDK, ma richiede un'azione specifica da parte dei Publisher.
		Photo Galleries	L'SDK ha la funzionalità di conteggiare i contenuti interni alla pagina ma non li conteggia automaticamente. Il conteggio di tali elementi dipende da una scelta dell'editore. Nielsen raccomanda di utilizzare questa funzionalità quando cambia almeno il 50% del contenuto ma resta prerogativa di Audicom definire tali regole.
	Time Spent	Granularity	Granularità Tempo Speso: tempo di visualizzazione conteggiato in secondi e riportato in UI in minuti (59 secondi saranno riportati come 0 minuti)
		Visibility / Focus	Browser e tab devono essere in focus

		Timeouts	Date le condizioni Visibility/Focus, 30 minuti di inattività chiudono la sessione e vengono accreditati alla pagina
Mobile App	Launches	Start Condition	Mobile App in Foreground
	App PageView	Start Condition	Mobile App in Foreground Il view ping è inviato solo per la Home Page dell'App
	Time Spent	Minimum Granularity	Granularità Tempo Speso: tempo di visualizzazione conteggiato in secondi e riportato in UI in minuti (59 secondi saranno riportati come 0 minuti)
		Focus	Solo una app mobile alla volta può essere considerata in foreground
		Timeout	Quando la App mobile va in background Nota: SDK invierà un nuovo App Launch ping quando la App mobile: viene riaperta – viene chiusa e riaperta in un intervallo di tempo superiore ai 5 minuti – viene portata in background e riportata in foreground in un intervallo di tempo superiore ai 5 minuti
		Mobile Interruptions	Le interruzioni sulla fruizione mobile (ad esempio chiamate, ricezione di messaggi) sono escluse dal Tempo Speso

5.2.2. SDK Text– certificazione

Nielsen lavora direttamente con i Publisher iscritti per implementare gli SDK Text e validare la loro corretta strumentazione ai fini della misurazione.

Gli step relativi al processo di certificazione di SDK Text sono riportati di seguito.

Prima che il player sia certificato, Nielsen (tramite Operations team) supporta gli sviluppatori verificando tool e ambiente di test disponibili presso il Publisher.

- I test vengono eseguiti registrandone gli output (tramite *charles proxy*)

- Input Test: controllo sugli eventi inviati da player a SDK Text
- Output Test: controllo sugli eventi registrati da SDK Text e verifica che tutto sia correttamente inviato agli ambienti di collection dei dati. Il team Operations invia il feedback al cliente ed eventualmente fornisce ulteriore supporto nello sviluppo tecnico
- Nel momento in cui la certificazione è superata viene inviata una documentazione al cliente

5.2.3. Attività di controllo a cura di terze parti

L'istrumentazione dell'SDK Text viene mensilmente sottoposta a procedure di verifica ad opera di terze parti indipendenti (PwC Business Services SRL), nell'ambito dell'incarico conferito alle stesse da Audicom.

5.2.4. SDK Text – processo di raccolta dati

Ogni qualvolta un device utente invia una richiesta e carica un contenuto taggato (browser o app), l'indirizzo IP viene incluso nel processo comunicativo tra l'SDK e i server di raccolta.

Anche a fini di compliance normativa, tra le altre cose gli indirizzi IP rilevati non vengono trasferiti ad alcun server al di fuori del territorio europeo, e vengono anonimizzati prima di qualsiasi trasferimento, indipendentemente dal Paese di destinazione.

L'indirizzo IP, prima di essere anonimizzato, viene sottoposto a geo-localizzazione a livello di città come massimo livello di dettaglio.

Il processo per rendere anonimi gli indirizzi IP avviene interamente sui server di raccolta.

Se un browser non accetta cookie o un cookie non può essere creato, l'indirizzo IP viene mascherato forzando a zero (v4) l'ultimo ottetto; questa informazione viene concatenata allo User Agent, e la stringa risultante viene sottoposta a hash per creare uno "pseudo-cookie". Un indirizzo IP troncato più lo User Agent può ricomprendere numerosi dispositivi (ad es. tutti i computer Apple e iPad che utilizzano Safari e a cui l'Internet Service Provider abbia fornito indirizzi IP che condividono i primi tre ottetti, quindi fino a 256 indirizzi); inoltre la funzione di hashing è irreversibile ed è impossibile risalire all'IP troncato e allo User Agent.

Tutti gli indirizzi IP sono quindi sovrascritti in memoria e trasmessi ai servizi di produzione a valle come 0.0.0.0, quindi l'indirizzo IP originale è effettivamente cancellato, con l'unica eccezione marginale dei log conservati per un periodo di 180 giorni nel server di raccolta per motivi di sicurezza, in accordo con l'attuale regolamentazione, separati dai dati di utilizzo. I log non vengono trasferiti fuori dal territorio Europeo.

Il trattamento dei dati qui descritto potrà subire variazioni dipendenti dall'implementazione di ID System, come descritto nei paragrafi successivi.

5.2.4.1. Filtro per il traffico non umano (Anti-Fraud)

La raccolta di dati censuari è soggetta all'inevitabile raccolta di attività robotiche, ovvero tutte quelle attività che non sono riconducibili a utenti umani. Si rende quindi necessario l'utilizzo di filtri in grado di identificare tale traffico nella misurazione delle audience. Questo principio vale per la misurazione del traffico proveniente da browser e applicazioni.

Vengono utilizzati due tipi di esclusioni:

- L'esclusione di indirizzi IP e User Agent: consiste in una blacklist di spider, web-crawler e altri robot automatici. Alla base dell'esclusione c'è una lista standard, redatta una volta al mese dall'Interactive Advertising Bureau in USA (IAB) e dall'Audit Bureau of Circulations nel regno Unito (ABCe). Questa lista include gli indirizzi IP e gli User Agent dei robot che sono noti e che vanno quindi esclusi.
- L'esclusione activity-based: specifica le soglie massime di attività che possono essere considerate umane. Per identificare queste soglie vengono anche utilizzate le informazioni provenienti dai panel desktop e mobile. Il traffico per specifici cookie o pseudo-cookie viene escluso dinamicamente quando l'attività a loro associata supera determinati valori nelle 24 ore.

5.3. “Audicom Daily/Weekly Digital”

L'approccio utilizzato include le seguenti fasi fondamentali:

1. Gli SDK implementati dagli editori che sottoscrivono “Audicom Digital” permettono di misurare ogni singola visualizzazione del contenuto; si tratta della cosiddetta “Misurazione Censuaria”, che rappresenta una delle basi fondamentali di questa metodologia. I conteggi censuari sono fondamentali per rilevare i volumi ma non forniscono informazioni sull'audience. Per assegnare informazioni relative all'età e al genere viene fatto affidamento sul Nielsen DCR ID System che, partendo da “Big Data” forniti da Data Provider esterni, genera profili individuali anonimi che vengono associati alla misurazione censuaria volumetrica.
2. I Data Provider forniscono informazioni funzionali alla attribuzione di caratteristiche demografiche (nello specifico sesso/età) dell'audience (da cui comunque Nielsen in nessun caso risalirà all'identità di alcuna persona fisica/soggetto interessato), ma bisogna sottolineare che non sono esenti da bias e imprecisioni tra cui, ad esempio, la possibile inaccuratezza demografica, misattribution dovuta alla condivisione dei devices tra membri della famiglia. La soluzione adottata per risolvere queste problematiche è l'utilizzo di fattori di calibrazione in grado di eliminare i bias contenuti nelle informazioni fornite dai Data Provider (questo processo viene chiamato “calibrazione”).
3. Una volta calibrata l'assegnazione dei profili e modellizzata la parte di informazioni non direttamente profilata, si calcola la Unique Audience sulla base delle frequenze delle impression. E, infine, viene generata la reportistica.

5.3.1. L'uso di *DCR ID* per la misurazione dell'audience

Il DCR ID è il “motore” utilizzato per l'assegnazione delle informazioni demografiche di genere ed età e la deduplicazione delle audience tra device e tra device e persone.

Il DCR ID è un insieme integrato di tutti gli assets Nielsen inclusi quelli forniti dai Data Provider (come ad esempio hashed email, Mobile AD IDs e Cookie), selezionati da Nielsen anche in base alle dichiarazioni dagli stessi fornite in tema di compliance privacy, in grado di collegare identificativi di device a individui per quegli utenti che hanno prestato il loro consenso all'utilizzo dei propri dati a scopo di misurazione. Tutti i Data Providers attualmente utilizzati da Nielsen appartengono alla categoria dei Data Aggregators o a quella dei Panel Digitali.

Queste aziende aggregano dati di terze parti da una varietà di diverse fonti offline e online per fornire dati qualitativamente rilevanti per la profilazione delle audience digitali.

Tali dati possono essere dati di prima parte e provenire da fonti offline, ad esempio da società di telecomunicazioni o di servizi online che raccolgono dati di registrazione dei propri utenti (sia basati sul browsing che app) oppure società che raccolgono tali dati dagli iscritti ai propri sondaggi. Nielsen riceve, con cadenza regolare (ogni mese, ogni settimana, etc.) dai propri Data Provider dei database contenenti una serie di “Identity Records” (da cui comunque Nielsen in nessun caso risalirà all'identità di alcuna persona fisica/soggetto interessato) che consistono in stringhe composte da Hashed Email, device IDs e informazioni riguardanti sesso ed età. Questo avviene in maniera totalmente indipendente da quello che attiene alla raccolta censuaria tramite SDK. Il collegamento tra dato censuario e profilo derivante dal Nielsen DCR ID System avviene a valle della raccolta ed è uno step specifico del processo produttivo. I Data Provider non hanno evidenza dei Device ID raccolti tramite SDK, in quanto non necessario per il processo produttivo sopra descritto.

5.3.2. ID Graph

Il cuore del DCR ID è l'ID Graph. Esso rende possibile, partendo dall'insieme degli identificativi di device appartenenti ad un singolo individuo, la deduplica delle esposizioni ai contenuti digitali passando da un livello di device a uno di profilo individuale anonimo. Permette, inoltre, l'assegnazione delle caratteristiche demografiche (da cui comunque Nielsen in nessun caso risalirà all'identità di alcuna persona fisica/soggetto interessato) a partire dalle informazioni raccolte per ogni singolo device ID.

I collegamenti tra hashed email e Mobile AD IDs/Cookie-IDs (Device IDs) forniti dai Data Providers, infatti, servono anche come link alle informazioni demografiche per ogni device rilevato.

Tali collegamenti tra hashed emails e MAIDs/Cookies opportunamente elaborate da un algoritmo di “Community detection” permettono di creare “communities” (o cluster) che vanno a rappresentare profili virtuali (da cui comunque Nielsen in nessun caso risalirà all'identità di alcuna persona fisica/soggetto interessato). A valle di queste elaborazioni, viene effettuato uno snapshot dell'elaborazione dell'algoritmo, in cui ogni riga rappresenta

un profilo virtuale (da cui comunque Nielsen in nessun caso risalirà all'identità di alcuna persona fisica/soggetto interessato).

Lo step successivo per ottenere la misurazione avviene con il collegamento tra i profili inclusi nello snapshot, quale output del ID Graph, e i dati censuari. A tale scopo vengono utilizzati gli identificativi (MAIDs e Cookie al momento) che vengono raccolti dagli SDK e che possono essere ritrovati come caratteristiche dei profili individuati tramite l'ID Graph.

A seguito di questo match, i dati vengono aggregati in classi di genere ed età (le stesse usate in fase di reporting) e sottoposti a step di correzione e calibrazione. Questi step sono fondamentali per garantire la neutralità del sistema di rilevazione rispetto ai dati di terza parte e la sua indipendenza. Per le stesse ragioni, l'attuale processo di selezione dei Data Provider da integrare nella metodologia di misurazione è descritto nel dettaglio di seguito. Resta inteso che tale processo di selezione potrà subire delle modifiche anche in ragione dell'evolversi della metodologia e delle esigenze e degli sviluppi che interessano il mercato, fermo restando che tali modifiche saranno comunicate ad Audicom tempestivamente, e fermo restando l'impegno, contrattualmente assunto da Nielsen nei confronti di Audicom, ad essere sempre conforme alla normativa (anche privacy) applicabile.

5.3.2.1. Matrice di correzione delle demografiche

Come illustrato nel precedente paragrafo, prima dell'introduzione nel sistema di misurazione ogni Data Provider viene analizzato da Nielsen per verificare la qualità dei dati demografici raccolti. Il risultato dell'analisi è un modello predittivo che individua i dati demografici inesatti e fornisce regole per la loro correzione. Lo scopo finale è quello di correggere errori conosciuti e verificabili presenti nei profili dei Data Provider, prima di usare questi dati demografici nella misurazione dei contenuti.

Il database degli utenti registrati del Data Provider viene sottoposto da Nielsen a un'attenta analisi tramite comparazione dei singoli record degli utenti con le informazioni del panel. Nello specifico, viene creata una matrice probabilistica derivante dalle tabelle di incrocio tra Digital Panel (truth set) e Data Provider in grado di predire l'inaccuratezza delle informazioni di registrazione degli utenti del Data Provider e di correggerle. I fattori di calibrazione per demografica, determinati dalla matrice probabilistica, sono applicati ai valori volumetrici (es.: impression e al tempo speso). In particolare, sono applicati a livello aggregato, per il dettaglio di granularità massima disponibile: Channel-Platform (PC-Mobile) -Tipo di contenuto (Text/Video). Per tutte le aggregazioni superiori, essendo metriche sommabili, viene effettuato un roll-up a partire dall'aggregazione più granulare corretta.

5.3.2.2. Matrice della condivisione

L'errore di attribuzione si verifica quando la visualizzazione di un contenuto statico o video è attribuita ad un certo utente loggato al Data Provider, mentre in realtà chi si trova di fronte al device è un altro utente/membro della famiglia. Questo errore si verifica solo quando lo stesso device viene usato da diversi membri dello stesso nucleo familiare. Questo tipo di comportamento è indagato tramite una probability survey, e le informazioni così dedotte vengono utilizzate per creare una matrice di condivisione dei device. Tale matrice è calcolata

per categoria e sottocategoria di siti o app, per ogni classe demografica riportata e per piattaforma (PC, Mobile).

5.3.2.3. Fattori di Correzione della Non-Coverage

Definiamo la copertura (Coverage) come la percentuale di impression a cui è possibile assegnare caratteristiche demografiche (età e genere) tramite match diretto con la profilazione del Nielsen DCR ID System, sulla totalità delle impression raccolte dall'SDK.

I fattori di correzione sono generati tramite l'utilizzo di un algoritmo di ottimizzazione alimentato dai Nielsen data assets e dalle distribuzioni derivanti dalle snapshot. Esso è calcolato per ogni demografica riportata, a livelli di granularità differenti a seconda della numerosità del sample utilizzato come base di calcolo.

Il livello di granularità più dettagliato per il calcolo del Coverage Correction Factor è Sub-brand/piattaforma. Qualora il numero di casi dello snapshot per tale livello non fosse sufficiente per garantire una stima affidabile, si considera come base di calcolo la sottocategoria a cui il Sub-brand appartiene. Se anche in questo caso la numerosità non fosse sufficiente, si osserva la categoria associata al Channel in questione. In ognuno dei casi citati sopra, il calcolo del fattore di correzione della coverage è declinato per piattaforma (PC, Mobile) e applicato al livello aggregato (per classe demografica), per il dettaglio di granularità massima disponibile: Channel-Platform (PC-Mobile)-Tipo di contenuto (Text/Video).

5.3.3. Deduplicazione

Un altro obiettivo chiave per la misurazione delle audience digitali è tenere adeguatamente in conto la loro duplicazione su più siti e su tipi diversi di piattaforma (mobile e desktop).

Una volta calibrate le impression per demografica e riscalate sul totale censuario, al livello di aggregazione più granulare (sub-brand, piattaforma, tipo di contenuto (text/video)) viene effettuato il roll up delle impression per ottenere le volumetriche ai livelli di aggregazione superiori (cioè per le aggregazioni A. - G.):

- A. Brand
- B. Brand by Content Type (Text/Video)
- C. Brand by Sub-brand
- D. Brand by Sub-brand by Content Type (Text/Video)
- E. Brand by Platform (Computer/Mobile)
- F. Brand by Sub-brand by Platform (Computer/Mobile)
- G. Brand by Platform (Computer/Mobile) by Content Type (Text/Video)
- H. Brand by Sub-brand by Platform (Computer/Mobile) by Content Type (Text/Video)
→ Livello in cui le impression vengono calibrate e riscalate, per demografica

A questo punto si procede con la stima della Unique Audience, anche in questo caso a livello aggregato, utilizzando il seguente approccio:

1. Per le aggregazioni G. e H., le impression calibrate e riscalate vengono divise per la frequenza derivante dallo snapshot matchato con le impression per la stessa aggregazione
2. Per le aggregazioni B., D., E., F., le relative UA vengono calcolate utilizzando le UA calcolate nello step precedente (più granulari) a cui viene applicato un fattore di deduplicazione ottenuto dallo snapshot. In particolare:
 - a. Le UA delle aggregazioni B. e E. vengono calcolate a partire dalla somma delle UA dell'aggregazione G. stimate nello step precedente a cui viene moltiplicato il fattore di deduplica calcolato dallo snapshot.
 - b. Le UA delle aggregazioni D. e F. vengono calcolate a partire dalla somma delle UA dell'aggregazione H. stimate nello step precedente a cui viene moltiplicato il fattore di deduplica calcolato dallo snapshot.
3. La UA dell'aggregazione A. viene calcolata a partire dalla somma delle UA dell'aggregazione B. stimate nello step precedente a cui viene moltiplicato il fattore di deduplica calcolato dallo snapshot.
4. La UA dell'aggregazione C. viene calcolata a partire dalla somma delle UA dell'aggregazione D. stimate nello step precedente a cui viene moltiplicato il fattore di deduplica calcolato dallo snapshot.

Il calcolo delle UA avviene a livello aggregato, per ogni entità e per ogni demografica. Per fare in modo che non esistano inconsistenze tra le UA delle diverse aggregazioni a causa di arrotondamenti nel calcolo e per essere sicuri che le UA stimate non eccedano i valori degli universi, si applicano delle regole di “capping”.

5.3.4. SDK VIDEO *(a cura di Auditel S.r.l.)*

Lo scopo del servizio è quello di fornire ad Audicom, attraverso l'installazione della tecnologia SDK video, dati di fruizione digitale da device digitali abilitati alla visione via protocollo IP relativi ai siti web e app mobile degli editori aderenti alla Ricerca Integrata Audicom.

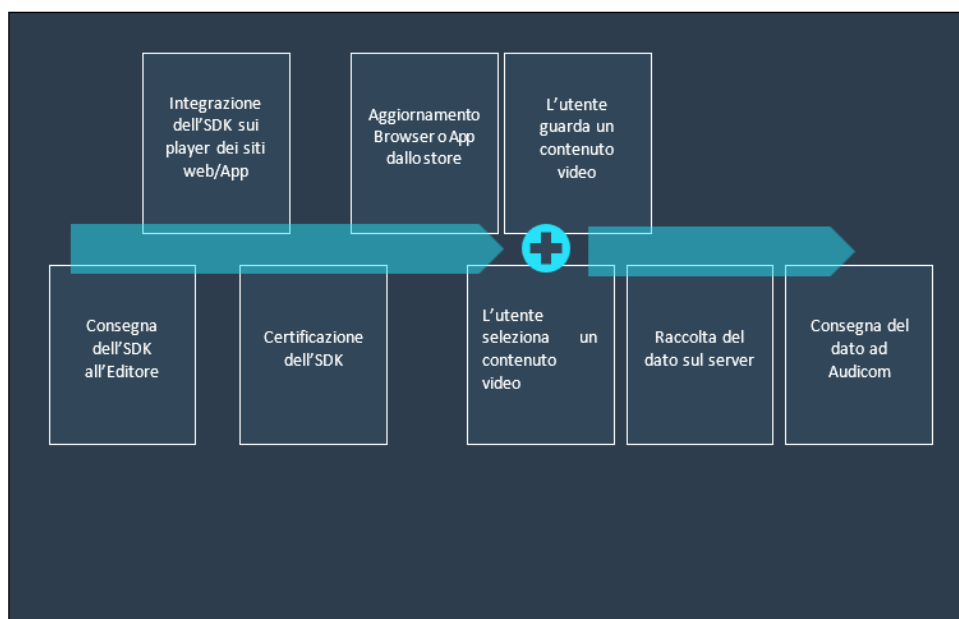
Il sistema rileva in modalità censuaria il consumo video riferito ai siti web e app mobile, aderenti alla Ricerca Integrata Audicom fruiti attraverso device digitali abilitati alla visione via protocollo IP (Tablet, PC, Smartphone, Set-Top-Box, Mini-Set-Top-Box, Game Console e CTV), grazie all'installazione della tecnologia SDK video direttamente nei player dei siti web e applicazioni mobile stessi.

5.3.4.1. Rilevazione Censuaria

L'SDK video viene implementato sulle piattaforme di visualizzazione su rete internet, ossia sui siti web e applicazioni mobile per Tablet, PC, Smartphone, Set-Top-Box, Mini-Set-Top-Box e Game Console aderenti alla Ricerca Integrata Audicom, al fine di rendere possibile la rilevazione in modalità censuaria, ossia la misurazione degli eventi di fruizione e visione (play, pausa, avanti veloce, stop ecc.) generati tramite i player su cui è lo stesso è integrato. A questi fini, l'SDK video raccoglie una serie di informazioni - l'orario e la tipologia di eventi (play, pausa, avanti veloce, stop ecc.), il tipo di web browser o mobile app utilizzati dall'utente e la tipologia di device - che consentono di rilevare in modo censuario i comportamenti di visione relativi ai siti web e app misurati, come nel seguito meglio dettagliato.

La tecnologia utilizzata richiede che l'implementazione dell'SDK video nel player avvenga nel rispetto di definite specifiche rispetto alle quale vengono svolte attività di supporto in fase di strumentazione, affinché l'implementazione avvenga in modo univoco e conforme ai protocolli, e attività di verifica una volta completata l'integrazione dell'SDK.

Lo schema sottostante mostra le varie fasi che portano alla messa in produzione di un player (ossia la effettiva rilevazione e produzione dei relativi dati).



i. Processo di instrumentazione del SDK e trasmissione dei dati generati dal Pairing Ping Nielsen

Come descritto precedentemente, la raccolta di dati avviene attraverso l'integrazione dell'SDK video nei player che fanno parte del perimetro di rilevazione della Ricerca Integrata Audicom.

Ai fini della Ricerca Integrata Audicom, vengono inoltre trasmesse tramite l'SDK Video informazioni generate tramite il Pairing Ping Nielsen che vengono da Auditel, se presenti e così come ricevute, trasmesse a Nielsen per la loro successiva elaborazione secondo quanto previsto dalla stessa Nielsen ai fini della Ricerca Integrata Audicom.

Una volta completato il processo di taggatura, Auditel avvia il processo di verifica della conformità dei player.

ii. Processo di certificazione del Player

Il processo di certificazione comporta l'esecuzione di test funzionali per verificare che l'implementazione dell'SDK sia conforme ai protocolli di instrumentazione.

Questo significa che:

- tecnicamente gli eventi (avvio, stop, pausa, fine, ecc.) vengono misurati correttamente;
- i metadati necessari sono correttamente valorizzati in base alle specifiche tecniche definite nell'apposita procedura.

I test vengono eseguiti su una serie di dispositivi rappresentativi di quelli normalmente a disposizione degli utenti (in particolare su una varietà di dispositivi allo scopo di realizzare una adeguata copertura del mercato). Nel caso in cui i test diano esito positivo e una volta che il player è stato portato in produzione dall'Editore sulla propria property, si procede a mettere a disposizione di Audicom i relativi dati censuari.

iii. Dispositivi e piattaforme rilevati

Vengono rilevati i contenuti video fruiti tramite web browser e mobile app su un'ampia gamma di dispositivi, di seguito elencati:

- Smartphone;
- Tablet;
- PC;
- CTV (sempre elaborato per la parte pubblicitaria, mentre il dato editoriale solo se sussistono le necessarie condizioni di produzione);
- Game Console;
- Dispositivi connessi (es. Google Chrome Cast, Apple TV, etc.).

iv. Tipologia di eventi rilevati

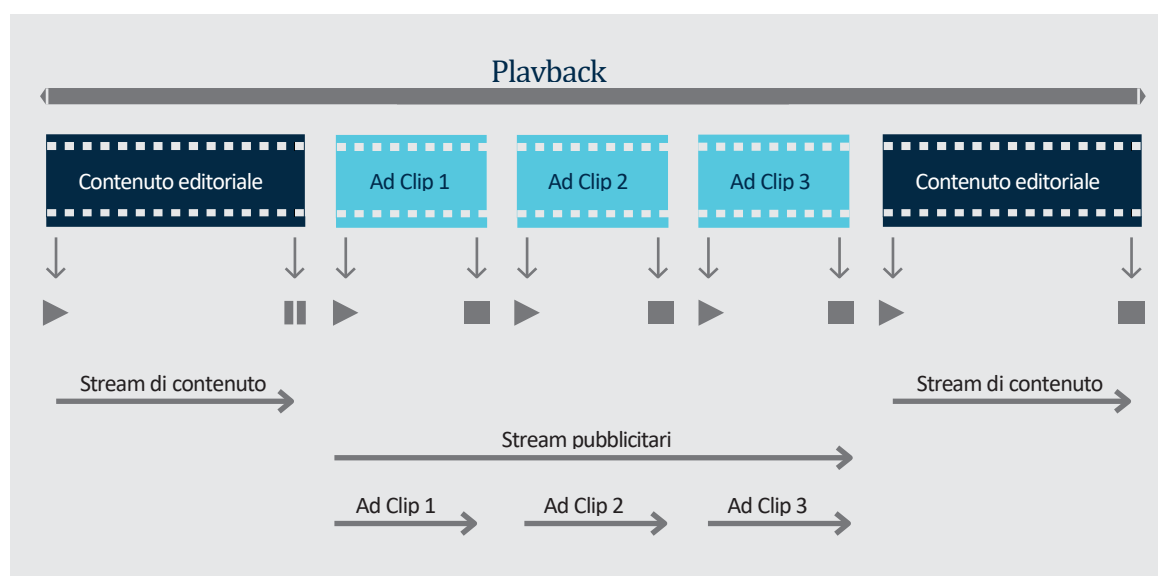
Tramite l'SDK video è possibile raccogliere una serie di eventi che caratterizzano la modalità di visione da parte dell'utente. La tabella sottostante riporta l'elenco di eventi che allo stato attuale vengono raccolti e processati per la produzione del dato. L'elenco potrà subire nel tempo variazioni e integrazioni in base all'evoluzione del progetto.

Lista di eventi	
Tipo di evento	Descrizione
Comando Play	Evento di inizio o ripresa della riproduzione di un video. L'evento può avvenire su espressa richiesta dell'utente o dovuto all'auto-play.
Comando Restart	Il video appena finito riparte dall'inizio. Tale evento può essere automatico o avvenire su espressa richiesta dell'utente.
Comando Pausa	Messa in pausa della riproduzione di in video.
Seeking	Navigazione all'interno del contenuto video.
Interruzione prematura del contenuto	Il video viene interrotto causa dell'abbandono della pagina, della chiusura del browser o dell'applicazione, ecc.

Pressione dell'home button (solo dispositivi mobili)	L'applicazione o il sito web vengono messi in background.
Pressione dello Sleep button (solo dispositivi mobili)	Il dispositivo viene messo in stand-by.
Buffering	Nel caso in cui la connessione non sia adeguatamente veloce, la riproduzione del video viene interrotta in attesa del caricamento della porzione di video non ancora riprodotta.
Alternanza contenuto editoriale e pubblicitario	Interruzione o la ripresa del contenuto editoriale per l'avvio o la fine di un'interruzione pubblicitaria

L'insieme di questi eventi permette di calcolare il tempo effettivamente speso e costituisce una caratteristica fondamentale della soluzione di misurazione SDK video.

Nel grafico sottostante viene riportato un esempio di comportamento di visione cui corrispondono gli eventi raccolti dall'SDK video.



5.3.4.2. Elaborazione dei dati e fornitura del dato censuario

Tutti i dati raccolti e generati dall'SDK video sono sottoposti a una serie di elaborazioni che consentono, tra l'altro, di eliminare la porzione di dati invalida (per esempio i dati generati da traffico web non umano) e produrre, così, i dati ai fini delle successive elaborazioni.

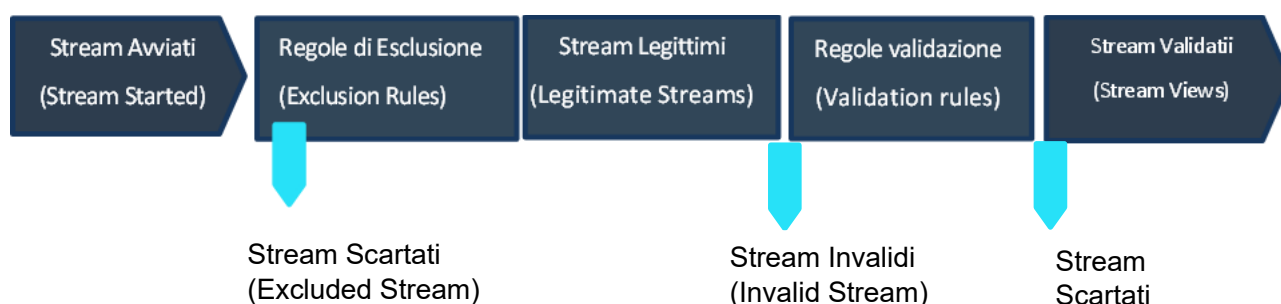
5.3.4.3. Elaborazione Dato Editoriale

A valle della raccolta dei Dati Editoriali, i dati che non provengono dal territorio italiano vengono eliminati e vengono elaborati solamente quelli realizzati all'interno del territorio italiano. Nell'elaborazione del Dato Editoriale non sono inclusi i dati CTV, contrariamente a quanto avviene per il Dato Pubblicitario.

5.3.4.4. Elaborazione Dato Pubblicitario

A differenza dei Dati Editoriali, i Dati Pubblicitari vengono considerati nella loro interezza, e non vengono pertanto filtrati in base all'area geografica. Allo stesso modo, i Dati Pubblicitari contengono anche il dato relativo alle CTV.

Il flusso di elaborazione prevede un processo di validazione che consiste in quattro fasi distinte:



Passo	Descrizione Passo
Stream Avviati (Stream Started)	Stream avviati. Sono tutti gli eventi (censuari) che vengono generati quando un video stream viene avviato su un dispositivo digitale.
Regole di Esclusione (Exclusion Rules)	Regole di esclusione. Si intende l'insieme delle regole Auditel che portano alla esclusione dei dati non idonei, ad esempio i dati risultati da traffico non umano. In particolare, relativamente alla esclusione del traffico non umano, la tecnologia utilizzata è in grado di individuare, filtrare e rimuovere il traffico non umano sia generico sia sofisticato (ovvero il traffico generato da chiamate in background che non richiedono malware sul device, come per esempio Spiders, Domain Laundering, Ad Injector)
Stream Scartati (Excluded Stream)	Stream scartati. Rappresentano gli stream scartati, in base alle regole di esclusione
Stream Legittimi (Legitimate Stream)	Stream legittimi. Rappresentano gli stream avviati che non vengono scartati dalle regole di esclusione.

Regole di Validazione (Validation Rules)	Regole di validazione. Si intende l'insieme delle regole che devono essere rispettate per considerare la fruizione di un video come effettivamente valida. Tali regole sono basate su aspetti oggettivi (per esempio possibili visioni simultanee di più contenuti sullo stesso dispositivo che ne generano lo scarto).
Stream Invalidi (Invalid Stream)	Stream invalidi. Rappresentano gli stream che vengono scartati in base alle regole di validazione. Questi stream non verranno presi pertanto in considerazione dal processo per il calcolo delle metriche.
Stream Validati (Stream Views)	Stream validati. Rappresentano gli stream avviati che hanno superato lo step di validazione

5.3.4.5. CUSV – Codice Unico Spot Video

E' un codice di tracciamento che viene aggiunto ai metadati del file video e alla descrizione delle sue caratteristiche lungo tutto il processo della sua erogazione verso la visualizzazione del consumatore: dal confezionamento del video nella fase di post-produzione, al caricamento sugli AdServer – delle concessionarie o dei centri media – alla erogazione nei player dei Broadcaster fino alla raccolta del codice di tracciamento da parte dell'SDK Video.

Lo stesso codice viene attribuito ai video erogati nei palinsesti della TV in modo che i Broadcaster ne tengano traccia nei loro sistemi e mettano in condizione Auditel di convertire i log file degli stessi Broadcaster aggiungendo il codice di tracciamento. Questa soluzione quindi consentirà di fare analisi cross TV-digital di una campagna pianificata sia in Tv sia sulle properties digitali (browser e APP) dei broadcaster.

- Il Codice Univoco Spot Video (CUSV) potrà misurare contenuti erogati direttamente o mediante Redirect su AdServer di terze parti.
- L'utilizzo del codice garantisce compatibilità con le varie tipologie di AdServer (Google, FreeWheel, Adagio, etc.) ovvero essere "agnostico" rispetto alla tecnologia di AdServer utilizzata.
- L'utilizzo del codice garantisce facilità di implementazione della soluzione da parte di coloro che devono caricare i contenuti sui vari AdServer.

5.3.5. Integrazione dei Dati Censuari Video *(a cura di The Nielsen Media Italy S.r.L.)*

Audicom utilizza un sistema che consente di utilizzare il dato censuario Auditel, raccolto tramite l'SDK Video Unico, all'interno del medesimo sistema di misurazione utilizzato per il traffico statico.

La soluzione è basata su uno scambio di *Log File* trasmessi con modalità *server-to-server* fornita da Nielsen, con l'aggiunta di una chiamata da operare a livello client per i browser da parte dell'editore.

Nello specifico, tale soluzione è in grado di gestire contemporaneamente i dati censuari relativi ai vari editori che abbiano implementato l'SDK Video Unico basandosi su:

- **Dati di visualizzazione:** *Log File* che contengono le informazioni relative ai volumi di consumo dei contenuti video raccolti ed inviati da Auditel ad intervalli regolari durante la giornata.
- **Ping di accoppiamento** - che gli editori implementano in modo che una chiamata parta al momento in cui l'utente fruisce di un contenuto video.

Il Ping di Accoppiamento è utilizzato come "gancio" per rendere possibile la stima delle audience dei dati volumetrici video attraverso il DCR ID, rendendo possibile l'utilizzo della stessa metodologia per la rilevazione dei contenuti Statici e Video.

5.3.6. Distribuzione Dati “Audicom Daily/Weekly Digital”

Audicom Media View: servizio multiplatforma riservato ai Soggetti iscritti che consente di accedere ai dati riferiti al sistema Audicom Daily-Weekly: i dati giorno e settimana si riferiscono alle Property di Publisher che siano stati certificati per tale tipologia di rilevazione; l'accesso all'interfaccia online ad essi dedicata viene consentito con password personalizzate fornite ai sottoscrittori del servizio.

I report saranno resi disponibili per le entità degli editori che abbiano provveduto a taggare le loro properties (sia tramite l'SDK di Testo che quello Video), tramite interfaccia Web comprensiva di dashboard interattive e di funzionalità per la creazione di report personalizzabili che, a loro volta, potranno essere scaricati in diversi formati o inviati via email. I dati saranno disponibili in forma aggregata per segmenti demografici per genere ed età e per i livelli gerarchici definiti come descritti nella sezione dedicata al catalogo.

La tabella seguente descrive in dettaglio le metriche disponibili per Audicom Daily / Weekly basato su Nielsen DCR ID e le relative formule di calcolo:

Metrica Primaria	Tipo di contenuto	Descrizione
App Launches	Video, Text, Video + Text	Il numero totale di volte in cui è stata aperta un'applicazione Precisione: numero intero
Page Views	Text, Video + Text	Conteggio delle singole pagine viste per uno specifico contenuto statico / testuale. Precisione: numero intero
Time Spent in Minutes	Video, Text, Video + Text	Totale del tempo speso, espresso in minuti, dalla totalità dell'audience Precisione: numero intero Video: Duration where content type = video Text: Duration where content type = text
Unique Audience	Video, Text, Video + Text	Il conteggio di individui unici esposti al contenuto.
Video Views	Video, Video + Text	Il numero totale di video starts. Ogni nuovo video viene considerato come un nuovo video start, sia in riproduzione automatica (autoplay) che avviato dall'utente Precisione: Numero intero

5.4. Privacy e Riservatezza

Nella massima misura possibile, la Ricerca Integrata utilizza dati anonimi e/o aggregati e, in ogni caso, i dati e le informazioni sono trattati unicamente ai fini di rilevazione audience, ossia di misurazione, in forma anonima e aggregata, della fruizione dei contenuti rilevati da parte degli utenti e sui volumi di consumo dei medesimi contenuti da parte degli utenti (e.g. contenuti fruiti, tempi di visione, tipologia di device nonché area geografica di riferimento).

Con particolare riferimento alla rilevazione tramite SDK, vengono raccolte e trattate informazioni che consentono di ottenere informazioni sui comportamenti di visione degli utenti (i.e. i contenuti consultati, la durata della visualizzazione, il tipo di dispositivo utilizzato e l'area geografica, gli indirizzi IP, le caratteristiche e gli identificatori del dispositivo utilizzato), ma non di risalire all'identità degli utenti. Sono altresì state implementate adeguate misure tecniche ed organizzative, conformemente alla normativa applicabile, affinché il potere identificativo degli SDK utilizzati per l'attività di rilevazione audience di cui alla Ricerca Integrata venga ulteriormente ridotto (e.g. con riferimento al mascheramento dell'indirizzo IP degli utenti) e le informazioni vengano anonimizzate nella massima misura possibile.

Audicom assicura l'osservanza della normativa privacy anche richiedendo ai propri fornitori l'assunzione di impegni espliciti e di garanzie specifiche in materia di conformità alla normativa applicabile – inclusa la normativa privacy.

6. AUDICOM CTV (a cura di The Nielsen Media Italy S.r.L.)

6.1. Obiettivo della rilevazione

La ricerca ha l'obiettivo di fornire dati relativi al consumo di contenuti video digitali attraverso le Connected TV. A partire dai dati di rilevazione effettuati da Auditel, il sistema è in grado di produrre stime di audience relative ai principali Broadcaster misurati da Auditel.

I dati sono resi disponibili mensilmente tramite il Database RL Digital.

6.2. Metodologia

L'approccio metodologico ha come obiettivo quello di elaborare, tramite modelli statistici, i dati CTV rilevati da Auditel, in modo da ricreare input simili a quelli che l'algoritmo sRLD riceve abitualmente dal sistema DCR ID. Questo consente di mantenere inalterato il dato Digital mensile, pur rendendo possibile espanderlo per includere una nuova piattaforma di visione.

La metodologia prevede l'utilizzo di modelli statistici specificatamente ideati e sviluppati per questa soluzione. Tali modelli permettono di:

1) stimare a livello probabilistico i panelisti digitali che hanno fruito dei contenuti di almeno uno dei Brand CTV nel mese in esame (Attivi CTV).

2) Calcolare le stime di reach, views e tempo speso, aggregate per:

- Brand CTV
- Caratteristiche demografiche
- Granularità temporale (es.: giorno, mese).

I dati aggregati così ottenuti dai modelli forniranno i constraint specifici, ovvero i target CTV, che l'algoritmo sRLD dovrà far convergere tra di loro, utilizzati nell'ultima fase del processo dell'algoritmo stesso.

In questo ultimo step produttivo ad ogni panelista, considerato "Attivo CTV", verranno quindi assegnati volumi di visione in maniera coerente per ognuna delle variabili. Verrà effettuata quindi la medesima applicazione dell'algoritmo sRLD utilizzato nell'attuale produzione. In questo modo viene assicurata la coerenza dell'elaborazione congiunta delle informazioni di esposizione digitale con quelle di esposizione a contenuti CTV.

6.3. Distribuzione del Dato

I dati relativi alla Piattaforma CTV saranno integrati nel Database RL Digital e pertanto seguiranno le caratteristiche di consegna, formato e metriche.

Il Database RL Digital manterrà il formato descritto nel relativo Allegato, gli specifici dati CTV saranno ottenibili tramite uno specifico codice del campo PIATTAFORMA all'interno del file AWDB_MMYT_STDFMT_NAVI_SRLDa.DAT.

6.4. Periodo di rilevazione

Rilevazione continuativa

6.5. Indirizzo web di pubblicazione della metodologia

www.audicom.net

7. AUDICOM DATABASE RESPONDENT LEVEL DIGITAL

7.1. Introduzione

Il principale risultato al quale si arriva con i dati raccolti dal panel Computer e Mobile è mettere a disposizione degli utilizzatori (pianificatori, concessionarie ed editori) una base di dati, Audicom Database Respondent Level (Digital) che possa essere elaborata con gli appositi software di Analisi delle Audience e di Valutazione dei Piani.

Tale base di dati deve essere a livello individuale (respondent level), cioè contenere, per i singoli individui del campione rappresentativo degli utenti del mezzo, i dati descrittivi dell'individuo in termini sociodemografici e di consumo del mezzo stesso.

Audicom Database Respondent Level Digital comprende siti e oggetti pianificabili di Property iscritte al sistema Audicom e Brand «Ad Supported» rilevati esclusivamente mediante Audicom Panel con Total Digital Audience, computer & mobile, media mensile pari o maggiore a 20 milioni e appartenenti a Parent con Total Digital Audience, computer & mobile, media mensile pari o maggiore a 30 milioni. (Periodo di valutazione della U.A.: rilevazione solo Panel, no SDK, Ott-Dic 2019). I dati sociodemografici e il livello di informazioni a base singoli individui sono necessari per permettere analisi e valutazioni di piani su target definiti liberamente dagli utenti.

La descrizione del consumo del mezzo da parte dei singoli individui è rappresentata come descritto in Allegato 10.1. e si riferisce a ciascuno dei siti/canali per le Property iscritte ad Audicom Digital ed è invece descritto con la denominazione “Altro” per tutti i siti/canali delle Property non iscritte (ad eccezione dei casi descritti qui sopra).

La costruzione del dato Audicom Database Respondent Level Digital mensile rilasciato alle Software House avviene secondo la procedura delineata nei paragrafi successivi. Esso conterrà inoltre i dati ottenuti tramite la rilevazione CTV.

7.2. Metodologia Synthetic Respondent Level (per il dato digitale)

La metodologia alla base del Synthetic Respondent Level Data (sRLD) replica i panelisti del Total Digital Panel (ottenuto dalla fusione dei panel online e mobile) e ottimizza il loro comportamento per bilanciare l'integrità del panel con le metriche target derivate dalla misurazione censuaria. Per raggiungere questo obiettivo è necessario suddividere i dati in tre categorie:

- *Content* – definisce i brand e i sub-brand per i quali viene generato il respondent level database
 - Brand: associazione dei codici dei brand ai rispettivi sub-brand
 - Sub-brand: associazione dei codici dei sub-brand ai rispettivi brand

- *Universe* – le caratteristiche dei panelisti vengono utilizzate per stimare l'universo di riferimento
- *Panel* – caratteristiche e comportamenti dei panelisti
 - Caratteristiche: informazioni demografiche da replicare
 - Comportamenti: contenuti visualizzati da bilanciare con i target censuari

La soluzione per la misurazione dei contenuti digitali attraverso DCR ID è stata sviluppata con riferimento all'ecosistema dei media digitali. Il Synthetic Record Level Data è un processo che permette di calibrare le misure rilevate nei panel rispetto al dato proveniente dal DCR ID attraverso un algoritmo che replica i panelisti e corregge le misure rilevate.

Tale processo avrà la prerogativa di:

- Riportare la navigazione per singolo individuo per “entità” partendo da un dato aggregato
- Garantire la presenza delle Demografiche presenti nel panel, oltre ad Age by Gender
- Mantenere la misurazione del mercato sia per contenuti taggati che non taggati
- Includere nel sRLD contenuti taggati che per dimensione non sono rilevabili dal solo panel
- Far quadrare il dato mensile sRLD con DCR ID System

7.2.1. Processo di Fusione dei panel digitali

Il processo di Fusione permette di “fondere” i diversi panel digitali, che sono pesati separatamente, in un unico database. Il processo è diviso in 4 fasi utili alla determinazione della Total Digital Audience (TDA).

Fase 1

Identificazione, nel campione Smartphone, della porzione “Mobile Only”, che rappresenta la popolazione attiva negli ultimi 30 giorni da mobile AND senza accesso ad internet da PC. Tale insieme è identificato attraverso la rilevazione delle informazioni raccolte nel questionario di reclutamento.

Fase 2

Fusione all'interno dell'insieme “Mobile Only”: i panelisti attivi Tablet AND senza accesso ad internet da PC (donatori) sono fusi con i panelisti attivi Smartphone AND senza accesso ad internet da PC (riceventi). In tale fase si assume che l'insieme Tablet è totalmente incluso nell'insieme Smartphone poiché l'insieme Tablet Only rappresenta una componente molto piccola dell'universo totale.

Fase 3

Fusione tra PC panel (ricevente) e panel Mobile (donatore), esclusa la porzione di “Mobile Only”. In questo caso si assume che l'insieme PC non contenga totalmente quello Mobile

garantendo così la gestione dell'insieme "Mobile Only" che rappresenta una novità del processo.

Fase 4

L'unione degli "Individui con accesso ad internet da PC con attività Mobile" (Fase 3) e degli "Individui attivi da Mobile AND senza accesso a internet da PC" (Fase 2) rappresentano la popolazione TDA schematizzata nella tabella successiva.

Alla popolazione TDA viene aggiunta la popolazione non-digital, che rappresenta l'universo di individui non attivi da Mobile AND senza accesso a internet da PC viene generata (stimata e pesata) utilizzando le informazioni provenienti dalla ricerca di base (establishment survey).

Popolazione TDA		
Computer	Smartphone	Tablet
x	x	x
x	x	
x		x
x		
	x	x
	x	

7.2.2. Soft Calibration

Il processo di Soft Calibration è una tecnica di pesatura utilizzata per ottimizzare l'allineamento tra il Panel e il DCR Identity DCR ID. L'obiettivo è quello di calcolare un peso iniziale, per ogni individuo appartenente all'insieme TDA definito nel processo della Fusione, al fine di rendere l'algoritmo di SRLD più efficiente ed efficace.

La Soft Calibration, quindi, ricalcola i pesi del Panel (dopo la Fusione) tenendo in considerazione i seguenti vincoli:

- Demografiche Standard (Utilizzando le variabili usate per la pesatura degli attuali panel)
- Entità (Brand, Sub-brands) a totale (non incrociate by demo/device etc.)
- Valori di TDA per regione e Age by Gender derivanti dalla ricerca di base

Il processo individua come target le misure da DCR ID System delle entità taggate e interviene modificando il peso dei panelisti. Il processo avrà effetto sui pesi dei panelisti, sia delle entità taggate che delle entità non taggate.

La convergenza del processo è "rilassata" per alcuni target, cioè l'obiettivo della Soft Calibration è quella di raggiungere i vincoli prefissati prevedendo una certa tolleranza.

7.2.3. Algoritmo SRLD

Il core del processo è rappresentato dall'algoritmo, basato sulla metodologia Synthetic Respondent Level Data che clona i panelisti del Total Digital Panel (ottenuto dalla fusione dei panel online e mobile) e ottimizza il loro comportamento per bilanciare l'integrità del panel con le metriche target derivate dalla misurazione censuaria.

L'algoritmo prende in considerazione vari aspetti metodologici:

- La rappresentatività della popolazione – il dato prodotto rappresenterà la popolazione PC 4+ e quella Mobile 13-74
- Le metriche che saranno incluse nel database finale sono le seguenti:
 - Unique Audience (UA)
 - Page Views / Video Views
 - Time Spent

Tali metriche sono le metriche considerate nell'algoritmo per sRLD.

- Content type – il database finale produrrà il contenuto "Text" e quello "Video" solo per le entità taggate
- Piattaforma – il database finale conterrà la navigazione da Browser per PC e la navigazione combined (Browser+App) per Mobile
- Entità taggate - Le entità taggate saranno in linea con il dato DCR basato sull'ID System
- Entità NON taggate - Le entità non taggate non saranno impattate dall'algoritmo sRLD, le misure di tali entità non entrano nelle fasi successive del processo e saranno riportate nel dato sRLD finale come da misurazione panel, senza ulteriori modifiche
- Vincoli – Le misure prodotte nel database RLD saranno funzione dei vincoli definiti (provenienti da DCR basato sull'ID System)
 - Popolazione: 4+ PC, 13+ Mobile
 - Segmenti Demografici: Gender by Age da DCR ID System
 - Entità: Brand, Sub-brand taggati (Parent calcolato a valle). Garantita la coerenza tra Sub-brand e Brand (Somma dei Sub-brand appartenenti al Brand x sarà uguale al Brand x da DCR ID System)
 - Time frame: monthly e daily DCR ID System. Sarà fornita anche la suddivisione in fasce orarie, non considerate come constraint.

Il Synthetic panel finale è determinato da una serie di parametri che sono stabiliti prima dell'avvio del processo e da una serie di loop che vengono eseguiti per ottenere una composizione ottimale.

Il primo passaggio prevede la clonazione dei panelisti e la distribuzione dei loro pesi tra i rispettivi cloni, ma a questo punto ciascuna copia possiede le stesse caratteristiche e gli stessi comportamenti del panelista di origine. In seguito, utilizzando questo nuovo panel “sintetico”, vengono calcolate le distanze dalle metriche finali desiderate. Queste distanze sono ciò che l’algoritmo andrà ad ottimizzare in maniera iterativa. Sia il numero delle copie dei panelisti che la distanza minima ritenuta accettabile rispetto ai vincoli sono configurabili.

Se il panel configurato risulta troppo distante dal raggiungimento dei vincoli desiderati, gli individui verranno selezionati sulla base dell’impatto che una loro modifica avrà sulla composizione del panel. Per esempio, un panelista che influenza molte metriche, alcune nel senso desiderato ed altre no, potrebbe non essere il candidato ideale da utilizzare per l’ottimizzazione del processo. Questa fase del processo è in parte calcolata tramite un ‘penalties set’ configurabile.

Infine, il comportamento dei panelisti selezionati viene modificato, o in aumento o in diminuzione/rimozione, e la distanza dalle metriche target viene nuovamente calcolata. Questo processo è poi ripetuto finché il panel di origine con i panelisti replicati (synthetic) contiene comportamenti idonei al raggiungimento delle metriche target desiderate.

Il risultato finale è un panel le cui stime di navigazione per le entità taggate corrispondono alle stime rilevate censuariamente. Le entità non taggate resteranno incluse nel database come da misurazione panel, senza essere modificate in alcun modo dall’algoritmo.

Il database finale includerà anche gli individui “non utilizzatori Internet” in misura coerente con le informazioni provenienti dalla Ricerca di Base fornita.

7.3. Rilevazione e gestione del traffico dei contenuti distribuiti tramite mobile In-App Browsing

Respondent Level

Il processo di gestione del traffico relativo a contenuti distribuiti attraverso mobile In-App Browsing avviene come di seguito descritto. In fase di rilevazione censuaria, per le entità dotate di SDK, il traffico fruito tramite in-App browsing viene individuato. L'algoritmo sRLD lo attribuisce alle diverse righe/respondent insieme al traffico organico del brand, sulla base del target DCR ID.

Reporting nella User Interface

Mobile In-App Browsing: il traffico relativo ai contenuti così distribuiti è incluso (=cumulato) nei report giornalieri e settimanali e Monthly dei Brand/Sub-Brand dell'editore. Lo spaccato con tale dato isolato non è disponibile.

Google AMP (per le entità per le quali è attivata la specifica rilevazione): il traffico è incluso (=cumulato) nei report giornalieri e settimanali e Monthly dei Brand/Sub-Brand dell'editore. Lo spaccato non è disponibile.

8. AUDICOM PRINT SURVEY

8.1. Quadro sintetico

8.1.1. Universo considerato

Popolazione italiana adulta (individui 13 anni ed oltre)

8.1.2. Campione per il rilevamento dei dati ed epoca del rilevamento

Indagine Single source

25.500 interviste all'anno: 10.000 casi CAPI e 15.500 casi CAWI (circa 13.350 adulti 18-64 anni, circa 1.350 Giovani 13-17 anni e 800 individui CSS 18 anni e oltre).

Ripartizione interviste in tre cicli di rilevazione all'anno, da circa 8.500 casi ciascuno (circa 3.330 interviste CAPI e circa 5.170 interviste CAWI), rilevati nell'arco di circa 32 settimane:

- primo ciclo: da fine gennaio a metà aprile
- secondo ciclo: da fine aprile a metà luglio
- terzo ciclo: da metà settembre a inizio dicembre

Avvertenza: le informazioni relative alla cumulazione dei cicli sono riportate nei paragrafi successivi

8.1.3. Numero dei comuni di campionamento

Indagine Single source

Circa 350 comuni per circa 417 CDC (Centri di Campionamento) per le interviste CAPI e circa 1.900 comuni per le interviste CAWI, in ciascuno dei tre cicli di rilevazione.

8.2. Sintesi metodologica

8.2.1. Premessa

L'indagine "Audicom Print" adotta un impianto single source, che consente la rilevazione della lettura di quotidiani e periodici nello stesso questionario, utilizzando interviste personali condotte attraverso il sistema di rilevazione CAPI (Computer Assisted Personal Interviewing) - doppio schermo, su un campione rappresentativo della popolazione adulta a livello nazionale e attraverso il sistema di rilevazione CAWI (Computer Assisted Web Interviewing).

Attraverso il questionario che rileva la currency multiplatforma, l'indagine "Audicom Print" esplora le modalità di lettura degli italiani partendo dalla testata per poi approfondire le diverse versioni usufruite (approccio brand first), rilevando contestualmente la lettura delle copie cartacee e delle copie digitali replica.

I dati vengono presentati secondo i target di "Lettori carta e/o replica" e "Lettori carta"; viene, inoltre, pubblicato anche il dato relativo al totale "Lettori replica".

L'adozione di un questionario volto a rilevare congiuntamente tutte le categorie di testate, caratterizza la configurazione di indagine.

In un anno standard vi sono tre cicli di rilevazione, indicativamente tra: metà gennaio-marzo; aprile-metà luglio; metà settembre - dicembre. La rilevazione viene quindi sospesa solo nei periodi estivo e natalizio: si tratta dunque di un'indagine continuativa.

La pubblicazione dei dati avviene cumulando i tre cicli più recenti sia per i Quotidiani sia per i Periodici.

Ogni edizione dell'indagine "Audicom Print" si compone, per ognuno dei cicli di rilevazione che la costituiscono, di tre fasi, descritte in dettaglio nelle pagine che seguono:

- Campionamento
- Rilevamento
- Elaborazione

8.2.2. Le stime dei lettori

L'indagine "Audicom Print" è un'indagine campionaria sulla popolazione adulta italiana, che ha lo scopo di stimare il numero dei lettori dei singoli quotidiani, supplementi illustrati di quotidiani e periodici partecipanti all'indagine e di descrivere i lettori stessi secondo molte caratteristiche sociodemografiche, per consentire la valutazione della loro attitudine a diventare destinatari utili ("target") della comunicazione veicolata attraverso le diverse testate.

Gli individui intervistati nell'indagine costituiscono un campione probabilistico, statisticamente rappresentativo di tutti i cittadini italiani, uomini e donne, dai 13 anni compiuti in su, che vivono nell'intero territorio nazionale, esclusi i cittadini che vivono all'estero.

Gli individui così selezionati vengono sottoposti ad una intervista face to face nelle loro abitazioni con un questionario adatto a rilevare con determinate tecniche l'eventuale avvenuta lettura da parte dell'intervistato di ciascuna delle testate quotidiane e periodiche partecipanti all'Indagine "Audicom Print", nei periodi precedenti l'intervista (p. es. negli ultimi 7 giorni, o negli ultimi 30 giorni, ecc.).

La tecnica dell'intervista prevede inoltre l'uso, quale stimolo per facilitare il riconoscimento delle testate ed il ricordo degli eventi di lettura, di uno schermo indipendente da quello del PC dell'intervistatore, sul quale appaiono i logo a colori delle testate partecipanti all'indagine.

Un questionario auto-compilato analogo a quello utilizzato nell'intervista face to face, contenente i loghi a colori delle testate rilevate, viene utilizzato per le interviste svolte per il target Giovani (individui fra i 13 e i 17 anni), per il sovracampione delle Classi Sociali Superiori (CSS) e per una quota del target di individui maggiorenni.

Grazie al requisito di rappresentatività del campione di intervistati rispetto a tutta la popolazione adulta italiana, che è nota dalle statistiche ufficiali nel suo ammontare, nella sua ripartizione geografica e nella sua composizione secondo i caratteri demografici (consentendo pertanto a posteriori i controlli e le ponderazioni necessarie per ovviare ad eventuali difetti di proporzionalità), i dati sulla lettura delle testate ottenuti dal campione possono essere estrapolati all'universo di riferimento, per ricavarne stime in migliaia di lettori di ogni testata.

Le stime di lettura fanno riferimento a definizioni convenzionali del lettore di una data testata.

Le principali definizioni di lettore adottate per i tre target "Lettori carta e/o replica", "Lettori carta" e "Lettori replica" sono le seguenti:

Per i QUOTIDIANI

- **"lettori (nel) giorno medio"**

- **Lettori "carta e/o replica"**

Quante persone leggono o sfogliano almeno una delle due versioni (copia cartacea, copia digitale replica) del quotidiano X almeno una volta, in media (nel corso di una settimana), in un giorno.

- **Lettori "carta"**

Quante persone leggono o sfogliano la versione cartacea del quotidiano X, almeno una volta, in media (nel corso di una settimana), in un giorno.

- **Lettori "replica "**

Quante persone leggono o sfogliano la versione digitale replica del quotidiano X, almeno una volta, in media (nel corso di una settimana), in un giorno.

Per i tre target si fa riferimento sia a persone che leggono il numero del giorno di lettura, sia a persone che leggono un numero arretrato; sia lettori che leggono una copia da essi stessi acquistata, oppure avuta in prestito, trovata, ecc. Non importa se la lettura ha avuto luogo in casa, o fuori casa, o in un luogo di lavoro, locale pubblico, ecc.

I lettori nel giorno medio vengono stimati calcolando la media delle persone che hanno letto il quotidiano in ciascuno degli ultimi 7 giorni.

Per i PERIODICI

- **"lettori ultimo periodo"**
- **Lettori "carta e/o replica"**
Quante persone leggono o sfogliano almeno una delle due versioni (copia cartacea, copia digitale replica) della testata X, almeno una volta, in 7 giorni se settimanale o supplemento settimanale di quotidiano, oppure in 30 giorni se mensile.
- **Lettori "carta"**
Quante persone leggono o sfogliano la versione cartacea della testata X, almeno una volta, in 7 giorni se settimanale o supplemento settimanale di quotidiano, oppure in 30 giorni se mensile.
- **Lettori "replica"**
Quante persone leggono o sfogliano la versione digitale replica della testata X, almeno una volta, in 7 giorni se settimanale o supplemento settimanale di quotidiano, oppure in 30 giorni se mensile.

Per i tre target si fa riferimento sia a persone che leggono il numero presente in edicola al momento della lettura, oppure un numero arretrato; sia a persone che leggono una copia da essi stessi acquistata, oppure avuta in prestito, trovata, ecc. Non importa se la lettura ha avuto luogo in casa, o fuori casa, o in un luogo di lavoro, locale pubblico, ecc.

La stima dei lettori nell'ultimo periodo, cioè delle persone che hanno letto un qualunque numero del settimanale X negli ultimi 7 giorni, o del mensile X negli ultimi 30 giorni, va considerata equivalente alla stima dei lettori del numero medio. Infatti, nell'ipotesi di lunga stabilità dei comportamenti di lettura del settimanale X, la stima delle persone che leggono in una data settimana, una o più volte, qualunque numero del settimanale X (il numero che è l'ultimo uscito e/o un numero arretrato), è equivalente alla stima delle persone che leggono in qualunque settimana (di un lungo periodo), una o più volte, un dato numero del settimanale X (sia che questo si trovi nella condizione di numero ultimo uscito, oppure in quella di numero arretrato).

In altre parole, sono "lettori ultimo periodo" tutti coloro che leggono o leggeranno, prima o poi, un dato numero di una testata, non solo nel periodo di presenza del numero nelle edicole, ma anche nei periodi successivi, se e fino a quando almeno una parte delle copie del numero sopravvivrà fisicamente nella disponibilità dei potenziali lettori.

Ulteriori dettagli sulle altre definizioni di lettura, sul tipo delle informazioni utilizzate per ciascuna stima e sul relativo metodo di rilevazione, sono riportati nella sezione intitolata "Rilevamento".

8.2.3. Campionamento

Universo rappresentato

La collettività cui si riferiscono i risultati dell'indagine ("Universo") è costituita dall'intera popolazione italiana adulta, così definita:

Tutti i cittadini italiani, di ambo i sessi, e in età dai 13 anni compiuti in su, che vivono nell'intero territorio nazionale, esclusi i cittadini che vivono all'estero.

Questa collettività viene stimata convenzionalmente in base alle statistiche ufficiali sull'ammontare della popolazione italiana residente e sulla sua distribuzione per regioni, sesso ed età, al lordo dei membri permanenti delle convivenze.

In particolare, per l'indagine "Audicom Print" i calcoli dell'ammontare della popolazione adulta per regioni, province e classi di ampiezza demografica dei comuni vengono eseguiti utilizzando le statistiche delle seguenti fonti ISTAT (ultimo aggiornamento disponibile):

- *Popolazione e movimento anagrafico dei comuni*
- *Popolazione residente).*

Metodo di campionamento

Il metodo adottato per la selezione del campione per le interviste CAPI si basa sui seguenti principi metodologici:

a) Utilizzo di tre stadi di selezione delle unità di campionamento (che sono: 1 - i comuni; 2 - le aree sub-comunali definite dalle sezioni elettorali; 3 - gli individui iscritti nelle liste elettorali sezionali).

b) Proporzionalità alla popolazione adulta italiana di una parte rilevante del campione.

c) Stratificazione (per il controllo della proporzionalità all'universo di riferimento) delle unità di campionamento secondo caratteri geografici (che sono: 1 - le province o gruppi di province (areole¹); 2 - le classi di ampiezza demografica dei comuni; 3 – la caratteristica capoluoghi di provincia / non capoluoghi²).

d) Sovracampionamenti (che comportano una deroga alla proporzionalità all'universo esposta al precedente punto b.), allo scopo di accrescere il grado di attendibilità delle stime nelle province o areole a dimensione campionaria critica, e cioè nelle province o areole in generale (nei casi in cui non raggiungano a totale e/o nelle celle di capoluogo o di fuori capoluogo una dimensione campionaria minima prestabilita) e in alcune province o areole

¹ Areola=singola provincia o gruppo di province nella stessa regione

² Si segnala che per l'indagine "Audicom Print" viene considerato un solo capoluogo per ciascuna delle province italiane esistenti. Per le province che nei file Istat hanno più di un comune capoluogo viene scelto come capoluogo la sede amministrativa della provincia.

in particolare, così da raggiungere la base campionaria minima per garantire l'accuratezza delle stime dei quotidiani nelle province che sono loro sede di edizione.

e) Casualità controllata della scelta delle unità di campionamento, assicurata mediante l'adozione di: 1) per i comuni - criterio sistematico casuale di estrazione su lista organizzata all'interno di ogni provincia, per tenere sotto controllo in modo ottimale la proporzionalità, sia secondo le sub-aree territoriali della provincia e sia secondo l'ampiezza demografica dei comuni; 2) per le sezioni elettorali delle grandi città – criterio sistematico casuale previa stratificazione secondo collegi elettorali; 3) per le sezioni elettorali dei comuni medi e piccoli - criterio sistematico casuale che utilizza la numerazione delle sezioni; 4) per gli individui da intervistare – criterio casuale dalle liste degli elettori iscritti nelle sezioni.

Il metodo adottato per la selezione del campione per le interviste CAWI si basa su un reclutamento di individui da access panel nel rispetto di una stratificazione puramente proporzionale alla popolazione adulta italiana rappresentata, per criteri territoriali e socio-demografici (campionamento per quota).

Maggiori dettagli sui metodi di campionamento vengono riportati nei paragrafi che seguono.

8.2.3.1. Metodo di campionamento CAPI

Stadi e unità di campionamento

Gli stadi del campionamento, come detto nell'enunciazione sintetica dei criteri metodologici, sono tre, ciascuno riguardante un determinato tipo di unità di campionamento:

1. Nel primo stadio vengono selezionate le unità primarie di campionamento, che sono i comuni, che vengono selezionati dall'universo di tutti i circa 8.000 comuni italiani, opportunamente organizzato per consentire le stratificazioni ed in generale tutti i controlli di rappresentatività e proporzionalità desiderati. Nel fare questa selezione vengono operati anche tutti i necessari sovracampionamenti, come sopra descritto (paragrafo “Metodo di Campionamento” punto d).
2. Nel secondo stadio vengono scelte, in ognuno dei comuni selezionati nel primo stadio, le unità secondarie di campionamento, che sono le aree sub-comunali, definite dalle sezioni elettorali.
3. Nel terzo stadio vengono scelte, in ognuna delle sezioni elettorali selezionate nel secondo stadio, le unità terziarie o finali di campionamento, cioè gli elettori, in quanto persone da intervistare senz'altro (adulti dai 18 anni in su), metà uomini, metà donne.

Sovracampionamenti territoriali

I sovracampionamenti, che costituiscono deroghe volute al criterio della proporzionalità geografica del campione all'universo, sono dettati dall'esigenza di potenziare determinate province o areole che, se trattate con il puro criterio della proporzionalità alla popolazione, risulterebbero coperte in misura insufficiente o non ottimale per rappresentare, entro margini di oscillazione accettabili, i fenomeni di lettura che interessano l'Indagine “Audicom

Print" (province o areole in generale di ampiezza demografica particolarmente ridotta e province sede di edizione di quotidiani locali).

I sovracampionamenti eseguiti sono descritti di seguito:

a. Sovracampionamenti provinciali minimali

Il "campione nazionale-base" viene distribuito fra le 71 areole (province o gruppi di province), in modo che alle celle campionarie dei comuni capoluogo e non capoluogo sia attribuita una quantità di interviste proporzionale alla popolazione adulta residente.

Nel caso in cui il totale di interviste attribuite alla provincia o all'areola /o al capoluogo o al non capoluogo della stessa sia inferiore a determinati valori soglia minimi prefissati (72 interviste per la provincia o areola e 24 interviste per il capoluogo e per il non capoluogo nel campione complessivo annuo), i totali sono aumentati sino ad ottenere i suddetti valori soglia.

b. Sovracampionamenti provinciali integrativi

Sono previste due ulteriori forme di sovracampionamento: 1) la prima consiste in un incremento generalizzato del 20% del campione base nelle province sedi di edizione dei quotidiani partecipanti all'Indagine "Audicom Print"; 2) la seconda consiste nel portare il numero minimo di interviste della provincia che sia sede di edizione di un quotidiano locale a 144 interviste e il numero minimo di interviste del suo capoluogo e del suo non capoluogo a 48 interviste nel campione complessivo annuo.

È importante ricordare che i sovracampionamenti descritti ai punti a, b non hanno come conseguenza necessaria la produzione di una stima di lettura più alta, ma assicurano soltanto una stima più accurata, poiché naturalmente le alterazioni della proporzionalità del campione rispetto all'universo, prodotte dal sovracampionamento, sono comunque eliminate (e pertanto la proporzionalità torna ad essere ripristinata) a posteriori mediante la ponderazione.

Campionamento dei Comuni

L'unità primaria (di primo stadio) del campionamento CAPI dell'indagine "Audicom Print" è data dal comune. Deve pertanto essere selezionato anzitutto un campione rappresentativo di tutti i comuni italiani. La fonte di questa selezione è data dagli archivi e relative statistiche dell'ISTAT, che sono disponibili a livello comunale.

Anzitutto viene calcolata la distribuzione delle interviste CAPI secondo province o areole e secondo comune capoluogo e altri comuni, con i criteri di proporzionalità e sovracampionamento voluti, basando l'attribuzione proporzionale delle interviste sulla popolazione adulta residente.

Le quantità di interviste così calcolate vengono poi divise per 8, essendo 8 l'unità modulare indivisibile di lavoro sul campo che è detta CDC = Centro di Campionamento. Il CDC è un

gruppo di 8 interviste da fare a 4 uomini e 4 donne, tutte in un breve intervallo di tempo, in un solo comune, e utilizzando nominativi estratti da 1 sezione elettorale del comune stesso.

Si ottiene in questo modo la distribuzione dei CDC tra le varie areole, e tra due tipi di comuni in ogni provincia o areola: capoluogo e non capoluogo (tutti gli altri). All'interno di questi ultimi vengono operate successive distribuzioni e ordinamenti dei comuni che consentono di tenere sotto controllo sia l'ampiezza demografica del comune e sia la sub-area territoriale nella quale il comune è ubicato.

La distribuzione ottenuta rappresenta il parametro-guida per la selezione dei comuni in cui effettuare le interviste.

La selezione dei comuni avviene operando sulla lista dei comuni di ogni provincia: lista che contiene tutti i comuni esistenti nella provincia, ad esclusione del comune capoluogo. La lista dei comuni non capoluogo di ogni provincia viene organizzata per sub-area territoriale della provincia e per ampiezza demografica, e cioè viene formata raggruppando i comuni prima secondo sub-aree nella provincia, poi ordinandoli nella sub-area secondo ampiezza demografica, ora crescente ed ora decrescente, e cioè facendo in modo che, passando da una sub-area all'altra, risultino adiacenti le stesse ampiezze demografiche.

Ai comuni capoluogo viene attribuito il numero di CDC previsto nel piano di distribuzione. In questo caso la scelta del CDC (zona/quartiere del capoluogo in cui effettuare le interviste) è organizzata nel quadro delle operazioni di estrazione delle unità di secondo stadio (sezioni), presso gli uffici elettorali dei comuni.

Per i comuni non capoluogo, la selezione dei comuni da campionare avviene nel modo seguente:

i comuni di ogni provincia vengono raggruppati per order (*variabile di stratificazione territoriale*) e ordinati per popolazione adulta all'interno dell'order (in modo alternativamente crescente e decrescente).

Al comune n-simo riportato nella lista di ogni provincia viene attribuito un valore corrispondente alla somma della popolazione del comune (P_n) più quella di tutti i comuni che lo precedono nella lista. Questo valore è detto valore della cumulata del comune n-simo (C_n). All'interno della provincia la popolazione totale di tutti i comuni (P_{tot}) viene divisa per il numero di CDC pianificati per quella provincia (c), ottenendo così il modulo di estrazione (M). Viene scelto casualmente un numero di partenza P che sia inferiore a M . Vengono calcolati c numeri di estrazione E , aggiungendo a P zero, una, due, tre, ecc. fino a $c - 1$ volte il valore M , e cioè: $E_1 = P$; $E_2 = P + M$; $E_3 = P + 2M$; $E_4 = P + 3M$; ...; $E_{(c-1)} = P + (c-2)M$; $E_c = P + (c-1)M$. Vengono scelti, come sedi di CDC, i comuni il cui valore di cumulata C_n è il più vicino ad uno dei numeri di estrazione E . Quando un comune ha una popolazione così ampia che il suo numero di cumulata C_n risulta il più vicino a più di uno dei numeri di estrazione, a quel comune viene assegnata più di una unità CDC.

Con questo metodo vengono assicurate le seguenti proprietà dell'insieme di CDC selezionati:

1. Una distribuzione nazionale dei CDC (e quindi delle interviste) equilibrata territorialmente, in quanto corrispondente alla distribuzione voluta del campione tra le province o le areole.
2. Una distribuzione dei CDC (e quindi delle interviste), all'interno delle province o delle areole, equilibrata per ampiezza demografica del comune e per ubicazione territoriale. Infatti: a) la probabilità di selezione di un comune (cioè di attribuzione di un CDC ad un comune) è direttamente proporzionale alla popolazione residente nel comune, poiché il margine ($C_n - C_{n-1} = P_n$) entro il quale può cadere un numero di estrazione E è ampio quanto la popolazione del comune; b) la probabilità di estrazione di un CDC appartenente a ciascuna delle sub-aree territoriali in cui è stata suddivisa la provincia è direttamente proporzionale alla popolazione residente nella sub-area.

L'effettiva sussistenza di queste 2 proprietà viene comunque garantita mediante un controllo a posteriori sul numero di CDC assegnato ad ogni classe di ampiezza demografica e ad ogni sub-area della provincia, con eventuale sostituzione di uno o più comuni, nei casi in cui l'attesa proporzionalità, per pura accidentalità, non si è esattamente verificata.

Le liste dei CDC estratti per le varie province o areole, parte capoluogo e parte non capoluogo, vengono cumulate in un'unica lista nazionale complessiva, in cui tutti i comuni e rispettivi CDC di tutte le province o areole sono disposti in un'unica successione, in cui ad ogni CDC è attribuito un codice di riconoscimento univoco, indispensabile per gli usi di lavoro delle fasi successive.

I comuni con più di un CDC compaiono nella lista tante volte quanti sono i CDC ad essi assegnati.

I CDC così estratti e "codificati" sono poi ulteriormente classificati ai fini del trattamento cui verranno sottoposti in fase di rilevazione, secondo i criteri seguenti:

- a. CDC in cui utilizzare una determinata edizione del questionario
- b. Settimana di rilevazione

Le suddette classificazioni sono realizzate garantendo una distribuzione casuale dei CDC, in osservanza al seguente criterio:

- c. Ripartizione casuale controllata per areola e comune capoluogo/non capoluogo dei CDC cui è assegnata una certa edizione del questionario.

La distribuzione casuale, controllata secondo i criteri sopra elencati, cerca di ottimizzare, compatibilmente con la minore dimensione dei sub-campioni rispetto a quella del campione complessivo, l'equivalenza della composizione geografica e quindi la confrontabilità

statistica tra tutti i sub-campioni, siano essi considerati per settimana di rilevazione o per diversa edizione del questionario.

Dopo la selezione dei comuni, come sopra descritta, può verificarsi ancora il caso che uno dei comuni selezionati non possa essere utilizzato per la rilevazione a causa dell'impedimento all'accesso alle liste dell'Ufficio Elettorale frapposto dalle Autorità Amministrative o per altri insuperabili impedimenti. L'Istituto in questi casi, dopo aver esperito anche direttamente (in sostegno all'azione dell'intervistatore) ma inutilmente ogni tentativo possibile di superare l'intralcio, chiede la sostituzione del comune.

Tale sostituzione è possibile nei casi in cui il comune interessato è abbastanza piccolo e viene scelto nella provincia o areola un altro comune adatto a sostituirlo, il più possibile all'interno della stessa classe di ampiezza e della stessa sub-area territoriale.

Campionamento delle Sezioni Elettorali

L'unità di campionamento CAPI di secondo stadio è data dalle sub-aree del comune, quali vengono definite dalle sezioni elettorali. Com'è noto, l'organizzazione della popolazione elettorale prevede l'aggregazione per sezioni, essendo ogni sezione formata da tutti gli elettori che risiedono entro una ben individuata e circoscritta sub-area del territorio del comune.

Per ogni CDC viene selezionata una delle sezioni elettorali del comune, e da ogni sezione vengono estratti i nominativi dai quali selezionare 9 intervistati (4 uomini e 4 donne, più un adulto da intervistare).

I metodi per il campionamento delle sezioni elettorali sono i seguenti.

Ad eccezione delle maggiori città, per le quali viene adottato un metodo speciale - la lista dei numeri di sezioni da estrarre viene predisposta nel modo seguente:

1. Il numero delle sezioni elettorali esistenti nel comune (N) viene diviso per il numero delle sezioni che devono essere estratte (n), che è pari al numero dei CDC. Questo quoziente è detto modulo di estrazione (M), e può essere anche un numero non intero.
2. Viene scelto casualmente un numero di partenza P , che non sia superiore a M . Il numero P corrisponde al numero S_1 della prima sezione da estrarre ($S_1 = P$).
3. I numeri delle rimanenti sezioni da estrarre ($S_2, S_3, \dots, S_n$) vengono scelti aggiungendo man mano M al numero P ($S_2 = S_1 + M$; $S_3 = S_2 + M$; ...; $S_n = S_{n-1} + M$).

Con il suddetto metodo la rappresentatività territoriale della popolazione compresa nelle sezioni estratte (nei confronti dell'intera popolazione del comune) è assicurata senza alcun rischio di errore accidentale soltanto nel caso che la numerazione delle sezioni risponda rigorosamente ad una logica di successione territoriale omogenea e continua (p. es. l'ordine a spirale, a partire dal centro, o altro ordine analogo). Nel caso che la numerazione delle

sezioni non segua un tale criterio, il metodo descritto assicura comunque la casualità (random) della scelta, anche se non assicura la protezione assoluta contro gli errori accidentali di campionamento.

Grandi città. Per le maggiori città viene adottato un metodo di selezione delle sezioni che assicura in modo assoluto (e cioè con la protezione anche contro gli errori accidentali di campionamento) la rappresentatività territoriale. In altre parole, il metodo assicura che la dislocazione delle sezioni estratte sul territorio comunale sia una riproduzione assolutamente fedele (proporzionale) di quella di tutte le sezioni, e quindi della dislocazione sul territorio di tutta la popolazione adulta del comune.

Per ognuna delle 7 maggiori città (Roma, Milano, Torino, Genova, Napoli, Palermo, Venezia) viene adottato il seguente metodo:

le sezioni del comune vengono suddivise secondo i collegi predisposti, fino al 2005, per la elezione della Camera dei Deputati. Com'è noto ciascun collegio corrisponde ad una determinata sub-area del comune. I collegi sono: 24 a Roma, 11 a Milano, 9 a Napoli, 8 a Torino, 6 a Genova, 6 a Palermo e 3 a Venezia.

I CDC assegnati al comune vengono suddivisi tra i vari collegi in proporzione alla popolazione elettorale (elettori iscritti). Per ogni collegio vengono estratte tante sezioni ordinarie quanti sono i CDC assegnati al collegio, più un adeguato numero di sezioni di riserva. Le sezioni ordinarie vengono scelte dopo aver disposto tutte le sezioni del collegio nell'ordine della numerazione delle sezioni, scegliendo ordinatamente una sezione ogni r , essendo r il rapporto tra S (numero delle sezioni esistenti nel collegio) e s (numero dei CDC, ovvero delle sezioni ordinarie da scegliere nell'area del collegio). Le sezioni di riserva (da utilizzare soltanto nel caso di impossibilità di utilizzare una sezione ordinaria) vengono scelte con lo stesso criterio. Si fa in modo di non scegliere come sezioni di riserva sezioni già estratte come ordinarie.

Dall'estrazione delle sezioni elettorali di un comune vengono escluse soltanto quelle corrispondenti interamente a comunità o convivenze (p. es. grandi ospedali).

Campionamento degli individui

La selezione delle unità di campionamento CAPI di terzo stadio, e cioè degli individui da intervistare, avviene in due momenti e ambienti distinti: 1) estrazione casuale di nominativi di elettori dalle liste delle sezioni estratte, negli uffici elettorali dei comuni selezionati; 2) scelta finale, fra i nominativi estratti, delle persone da intervistare, fatta dall'intervistatore, sul campo.

a. Scelta dalle liste elettorali

Di ogni sezione selezionata viene estratto un numero di nominativi pari al numero di interviste da fare (che per ogni sezione sono normalmente 8, di cui 4 a maschi e 4 a femmine). Viene anche estratto, per ogni nominativo ordinario, un adeguato numero di

nominativi di riserva. All'incaricato dell'estrazione, l'Istituto assegna i numeri d'ordine dei nominativi da estrarre da ciascun registro, che a loro volta sono stati estratti casualmente, nonché una serie di regole per la sostituzione del nominativo nel caso che il numero assegnato non possa essere utilizzato (perché il nominativo è depennato, o comunque non presente o perché il numero assegnato supera la numerazione occupata del registro). L'incaricato trascrive su moduli appositi i nominativi estratti, con l'indirizzo e l'anno di nascita, indicando anche il numero del nominativo nel registro e il numero della sezione.

b. Scelta finale sul campo

Sulle liste dei nominativi assegnati agli intervistatori, un quantitativo di nominativi pari al numero di adulti di 18 anni e oltre che devono essere intervistati è contrassegnato con la lettera "O" ("Ordinario"). L'intervistatore viene incaricato di intervistare proprio le persone corrispondenti ai nominativi contrassegnati con "O" (che sono destinati a formare il campione degli adulti di 18 anni e oltre).

Se un nominativo "O" non può essere intervistato per un motivo indipendente dalla volontà dell'intervistatore, questi lo sostituisce con un nominativo di riserva, scegliendo tra quelli dello stesso sesso che sono stati estratti come riserve. Fa eccezione all'uso del nominativo di riserva soltanto il caso della persona trasferita, che viene sostituita da un componente di almeno 18 anni della famiglia subentrata nell'abitazione, che sia dello stesso sesso e abbia l'età più vicina all'età del nominativo da sostituire. Se nessuna famiglia è subentrata oppure se nella famiglia subentrata nessun componente ha i requisiti richiesti, la sostituzione viene fatta, anche nel caso di trasferimento, ricorrendo ad un nominativo di riserva. Non è ammessa in alcun caso la sostituzione con una persona appartenente alla famiglia del nominativo che deve essere sostituito.

8.2.3.2. Metodo di campionamento CAWI

Campionamento Proporzionale

Per la parte di interviste online, la selezione degli individui da intervistare avviene attraverso access panel, in base a un metodo per quota.

In particolare, vengono studiati campioni proporzionali ad hoc per i seguenti target di popolazione:

- individui dai 18 ai 64 anni (target rilevato anche nel campione CAPI)
- individui Giovani dai 13 ai 17 anni

Sovracampionamento CSS

La parte di interviste online prevede anche un sovracampionamento del segmento di popolazione definito come "CSS = classi sociali superiori", cioè "imprenditori, dirigenti, liberi professionisti e loro familiari" allo scopo di accrescere il grado di attendibilità delle stime di

lettura, in quanto si presume che tale segmento della popolazione sia particolarmente esposto al fenomeno della lettura.

Per il sovracampionamento di CSS sono previste 800 interviste all'anno (circa 267 interviste CAWI per ciascun ciclo di rilevazione). Il sovracampionamento CSS viene selezionato con un metodo per quota.

Come per i sovracampionamenti territoriali CAPI, anche il sovracampionamento CSS non ha come conseguenza necessaria la produzione di una stima di lettura più alta, ma assicura soltanto una stima più accurata, poiché naturalmente le alterazioni della proporzionalità del campione rispetto all'universo, prodotte dal sovracampionamento stesso, sono comunque eliminate (e pertanto la proporzionalità torna ad essere ripristinata) a posteriori mediante la ponderazione.

Criteri per la distribuzione delle interviste CAWI

Per ciascuno dei tre target rilevati online, la distribuzione delle interviste è determinata come segue.

1. individui 18-64 anni: proporzionalmente all'universo su tutto il territorio nazionale, secondo i parametri di classificazione genere per area geografica, fasce di età per area geografica, regioni, ampiezza demografica dei centri della persona da intervistare
2. individui 13-17 anni: proporzionalmente all'universo su tutto il territorio nazionale, secondo i parametri di classificazione genere e area geografica del giovane da intervistare
3. segmento di popolazione CSS (Classi Sociali Superiori): con criterio proporzionale, a partire dalle statistiche derivanti dalla ponderazione preliminare ad hoc del campione complessivo Print 13+. La distribuzione viene controllata secondo i parametri di classificazione genere, fasce di età e area geografica della persona da intervistare. L'individuazione dei soggetti eleggibili avviene secondo la seguente definizione: la persona deve svolgere personalmente una professione eleggibile o appartenere ad una famiglia in cui un qualsiasi membro svolge o svolgeva prima della pensione una professione CSS.

8.2.4. Rilevamento

Attività di rilevamento

L'attività di rilevamento dell'Indagine "Audicom Print" consiste nella realizzazione di interviste personali e di interviste online.

Le interviste personali vengono effettuate nelle località prescelte (comuni e CDC), nei tempi (settimane) programmati e alle persone estratte a tale scopo dalle liste elettorali. L'Istituto ingaggia per tale attività molte centinaia di intervistatori free-lance, che per lo più risiedono nella zona in cui devono effettuare le interviste (nella stessa regione o provincia, se non

proprio nello stesso comune). Gli intervistatori vengono accuratamente addestrati prima delle interviste e il loro lavoro viene minuziosamente controllato subito dopo.

Queste interviste vengono effettuate face to face, nell'abitazione dell'intervistato o in luoghi adiacenti all'abitazione, seguendo un questionario completamente strutturato, con ben precisate modalità esecutive: vengono effettuate utilizzando un computer portatile attrezzato con programma speciale per la visualizzazione e gestione dei CDC, delle interviste e dei questionari assegnati, la registrazione assistita e controllata delle risposte e la trasmissione telematica dei dati delle interviste all'Istituto (sistema CAPI). Viene inoltre utilizzato un "tablet", cioè uno schermo sul quale compaiono i logo a colori delle testate partecipanti all'indagine.

Le interviste online vengono effettuate per una quota del target di individui maggiorenni, per il target Giovani (individui fra i 13 e i 17 anni) e per il sovracampione CSS. . Queste interviste vengono somministrate con un questionario auto-compilato attraverso il sistema CAWI. Il link al questionario viene inviato via e-mail ai rispondenti, su tutto il territorio nazionale. Nel corso dell'intervista il sistema CAWI permette all'intervistato di visualizzare i logo a colori delle testate rilevate e consente la registrazione assistita e controllata delle risposte, nonché l'invio telematico dei dati delle interviste all'Istituto.

Oggetto e strumenti del rilevamento

Oggetto del rilevamento dell'Indagine "Audicom Print" sono le singole letture, e cioè gli atti di lettura di singole testate quotidiane e periodiche, che vengono riferiti dagli intervistati agli intervistatori.

Per fare in modo che gli atti di lettura possano essere rilevati con riferimento ad un gran numero di testate, nel modo più completo possibile (senza dimenticare testate o eventi di lettura) e nel modo più preciso e coerente con il significato delle stime volute (in base a definizioni adatte degli eventi e dei tempi di lettura), vengono adottate nelle interviste determinate forme delle domande e procedure di somministrazione, che sono state verificate nella loro idoneità a consentire stime attendibili, esenti da problemi e anomalie di rilevamento.

In particolare, per consentire all'intervistato il riconoscimento delle testate da lui lette e per facilitargli il richiamo alla memoria degli atti di lettura avvenuti in determinati periodi di tempo, vengono presentate all'intervistato le riproduzioni in fac-simile del logo delle testate partecipanti all'indagine.

Per la rilevazione face to face CAPI tali riproduzioni, a colori, vengono visualizzate su "tablet": uno schermo, collegato al PC portatile dell'intervistatore mediante tecnologia Wi-Fi, cioè senza fili.

In fase di screening i logo delle testate visualizzati su tablet misurano circa 800 X 358 pixel e appaiono all'intervistato in modo "random", cioè in ordine sparso casuale, per evitare ogni rischio di possibile errore dipendente dal privilegio di posizione.

L'apparizione del logo su tablet è guidata dal programma installato sul PC portatile dell'intervistatore.

Per la rilevazione online CAWI, i logo a colori vengono visualizzati direttamente sullo schermo dell'intervistato insieme alle domande del questionario. La compilazione del questionario CAWI può avvenire da PC, da Tablet o, come alternativa secondaria, da smartphone. Anche nella rilevazione online i logo delle testate appaiono all'intervistato in modo "random" e la visualizzazione è guidata dal programma CAWI.

Questionario e metodo del rilevamento

Il rilevamento delle letture durante l'intervista avviene mediante un questionario che rileva congiuntamente la lettura di tutte le categorie di testata: quotidiani, supplementi, settimanali e mensili aderenti all'indagine "Audicom Print" per la rilevazione in atto.

Nel questionario è presente una sezione relativa alla visita dei siti web delle testate, derivata da domande somministrate esclusivamente a coloro che hanno dichiarato di aver letto la testata X su carta e/o replica (duplicazione della lettura con la visita del sito web della testata corrispondente, indipendentemente dalle modalità di accesso/fruizione del sito, ad esempio: accesso diretto al sito web della testata, ovvero al portale del gruppo editoriale, ovvero accesso attraverso un motore di ricerca, ecc.)

Le domande relative alla visita dei siti web delle testate dichiarate sono collocate al termine della sezione relativa alla lettura della testata su supporto cartaceo e/o digitale replica.

L'intervista si divide in tre fasi distinte tra loro. La seguente descrizione si riferisce all'intervista face to face CAPI, ma tutte le informazioni indicate sono rilevate anche nell'intervista online CAWI autocompilata.

Fase di screening (o di filtro):

Dopo l'introduzione, inizia la sezione di screening, che ha lo scopo di individuare e selezionare le testate lette o sfogliate almeno una volta, non importa in quale versione, cioè lette o sfogliate su carta o in versione digitale o quelle di cui è stato visitato il sito internet.

Per ogni testata viene perciò rilevata in un'unica domanda la lettura della testata nel suo insieme (versione cartacea, versione digitale uguale al formato cartaceo, versione digitale diversa dal formato cartaceo, sito internet). Il periodo di riferimento sono gli «ultimi 3 mesi» per i quotidiani e per le testate settimanali, gli «ultimi 12 mesi» per le testate mensili.

In questa fase viene fatta la selezione di tutte le testate che l'intervistato ha letto o sfogliato entro un periodo precedente abbastanza lungo, che è detto appunto periodo lungo, e che

è: 3 mesi per i quotidiani, per i supplementi settimanali di quotidiani, per i settimanali, e 12 mesi per i mensili; per ogni testata letta nel periodo lungo viene chiesto all'intervistato se, sempre in tale periodo, ha visitato anche il relativo sito Internet (se la testata ha un sito Internet).

L'intervistatore chiede all'intervistato di guardare i logo delle testate che si presentano, in ordine casuale, su tablet. L'intervistatore legge sul PC portatile il titolo della testata che in quel momento appare e chiede all'intervistato se l'ha letta o sfogliata almeno una volta nell'ultimo periodo lungo (indipendentemente dalla versione, dalla provenienza della copia e dalla data di uscita del numero). Se l'intervistato dichiara di avere letto o sfogliato la versione cartacea e/o versione digitale replica della testata almeno una volta nel periodo lungo, l'intervistatore chiede anche se ha visitato almeno una volta, sempre nel periodo lungo, il relativo sito web. Questa domanda viene chiesta solo se la testata ha un sito Internet.

Fase di pre-sviluppo:

Questa fase viene proposta in automatico dal software sempre e solo per le testate selezionate nella fase di screening come lette o sfogliate nel periodo lungo.

Le categorie di testate vengono trattate nell'ordine prescritto dal questionario, che segue il criterio qui sotto specificato a proposito dell'edizione del questionario.

L'ordine di presentazione di tutte le singole testate è casuale ed imposto dal programma.

La fase di pre-sviluppo riguarda la lettura delle testate selezionate dall'intervistato nella fase precedente. In particolare, viene chiesto all'intervistato di indicare quali versioni ha letto tra le diverse disponibili (su tablet appariranno solo le versioni esistenti e solo i dispositivi digitali su cui sono usufruibili). Pertanto, al fine di rendere più immediata l'individuazione della versione letta, è previsto per ogni testata un cartellino esemplificativo dei vari supporti su cui possono essere lette le diverse versioni di ciascuna testata.

1. La versione cartacea
2. La versione digitale uguale al formato cartaceo con eventuali contenuti aggiuntivi, che si può leggere o sfogliare dopo averla scaricata tramite app da tablet, da smartphone o da PC, da MAC
3. La versione digitale diversa dal formato cartaceo, scaricabile solo tramite app da tablet e/o da smartphone
4. Il sito internet

Per tutte le testate si chiude sempre ricordando: «La pagina che vede è **solo un esempio**, pertanto nel rispondere **non** faccia riferimento a quello specifico numero».

Nell'indagine sono presenti anche le domande sulla visita dei siti Internet delle testate. Queste domande vengono proposte in automatico dal software sempre e solo se l'intervistato ha dichiarato di aver visitato il relativo sito Internet nel periodo lungo nella fase di screening.

Nell'indagine vi sono 4 edizioni di questionario. Nell'edizione 1 l'ordine è il seguente: 1) Quotidiani, 2) Supplementi di quotidiani, 3) Mensili, 4) Settimanali. Nell'edizione 2 l'ordine è il seguente: 1) Quotidiani, 2) Supplementi di quotidiani, 3) Settimanali, 4) Mensili. Nell'edizione 3 l'ordine è il seguente: 1) Mensili, 2) Settimanali, 3) Quotidiani, 4) Supplementi di quotidiani. Nell'edizione 4 l'ordine è il seguente: 1) Settimanali, 2) Mensili, 3) Quotidiani, 4) Supplementi di quotidiani.

L'edizione del questionario ruota da un CDC all'altro.

Nella rilevazione *online* l'edizione del questionario ruota da un'intervista all'altra.

Fase di sviluppo:

Questa fase viene proposta in automatico dal software sempre e solo per le testate selezionate nella fase di screening e per le versioni selezionate nella fase di pre-sviluppo come lette o sfogliate nel periodo lungo.

Per ogni testata si rivolgono domande sulla frequenza di lettura; sull'epoca dell'ultima lettura e se l'ultima lettura ha avuto luogo nell'ultimo periodo (ultimo giorno, compreso negli ultimi 7 per i quotidiani, ultimi 7 giorni per i supplementi settimanali di quotidiani e per i settimanali, ultimi 30 giorni per i mensili).

Nell'indagine si chiede anche quante volte ha letto o sfogliato, entro quel periodo, una copia della testata. Per i quotidiani si chiede anche in quali degli ultimi singoli 7 giorni ha letto o sfogliato la testata.

Per i periodici si chiede anche quante pagine ha letto o sfogliato.

Nell'indagine si chiede anche qual è stato il modo di acquisizione (fonte di provenienza) della copia letta o sfogliata l'ultima volta nell'ultimo periodo.

Per quanto riguarda le domande sulla visita dei siti Internet delle testate, per i quotidiani si chiede la frequenza di visita del sito Internet nel periodo lungo, quando è stata l'ultima visita e se l'ultima visita ha avuto luogo nell'ultimo periodo (ultimo giorno compreso negli ultimi 7); si chiede, inoltre, in quali degli ultimi singoli 7 giorni ha visitato il sito Internet e quante volte, nell'ultimo giorno, è stato visitato. Per i periodici e per i supplementi di quotidiani si chiede se l'ultima visita ha avuto luogo nell'ultimo periodo (ultimi 7 giorni per i settimanali e i supplementi settimanali e ultimi 30 giorni per i mensili), la frequenza di visita del sito nell'ultimo periodo e quante volte, nell'ultimo giorno, è stato visitato.

Nel questionario, le domande sulla lettura di quotidiani e di periodici vengono rivolte in apertura di intervista, e sono seguite da altre domande, che hanno lo scopo di rilevare i parametri per la classificazione socio-demografica dell'intervistato, e cioè dati sull'intervistato e la sua famiglia: numero dei componenti della famiglia; sesso, età e condizione professionale dell'intervistato e di ciascuno degli altri componenti della famiglia; individuazione tra questi del "capofamiglia" e della persona "responsabile degli acquisti per la famiglia" (secondo la definizione tradizionale); grado di istruzione e titolo di studio dell'intervistato; stima fatta dall'intervistatore della classe socioeconomica e della classe di reddito mensile complessivo della famiglia.

8.2.5. Stime e definizioni di lettura

Tra le domande citate nel paragrafo precedente, alcune servono alla rilevazione delle informazioni necessarie per la più importante delle stime di lettura dell'indagine "Audicom Print", quella dei lettori ultimo periodo o dei lettori giorno medio (a cui si è già accennato nel paragrafo intitolato Le stime dei lettori) e sono quelle che, partendo dalla domanda di apertura rivolta nella fase di screening sulla lettura nel periodo lungo e arrivando a quelle rivolte nella fase di sviluppo e riguardanti i tempi dell'ultima occasione di lettura (o delle ultime), permettono di selezionare i lettori nell'ultimo periodo (da designare come quello breve, quando si vuole distinguerlo dal periodo lungo).

Quelle e altre domande rivolte nella fase di sviluppo forniscono invece le informazioni che, opportunamente trattate anche in combinazione tra loro, sono destinate a produrre dati strumentali (per controllo e calcolo) o integrativi (per corredo accessorio), come i dati sulla frequenza e probabilità di lettura per tutte le categorie; i dati sui lettori nei singoli giorni della settimana per i quotidiani (che peraltro vengono utilizzati anche per la stima dei lettori nel giorno medio); i dati sui contatti/copia nel giorno medio (per i quotidiani) o sui contatti/pagina nel periodo medio di lettura (per i periodici).

Riassumiamo qui di seguito le definizioni statistiche di lettura che vengono utilizzate per i dati pubblicati per i tre target "Lettori carta e/o replica", "Lettori carta" e "Lettori replica".

In particolare, per il target "carta" si fa riferimento a quanto rilevato dal questionario per la versione cartacea.

Per il target "carta e/o replica" le informazioni di lettura, per i lettori di entrambe le versioni, sono frutto di un ricalcolo finalizzato a combinare le informazioni rilevate separatamente dal questionario per la versione cartacea e la versione digitale replica.

Per il target "replica" si fa riferimento a quanto rilevato dal questionario per la versione digitale uguale al formato cartaceo. Per questo target sono presentati solo i dati relativi alle definizioni di lettori:

- nel giorno medio, per quanto riguarda i Quotidiani

- nell'ultimo periodo, per quanto riguarda i Supplementi di Quotidiani, i Settimanali e i Mensili

Per i QUOTIDIANI

a. Definizione-tetto: Lettori totali

Quanti leggono o sfogliano il quotidiano X (non importa quale numero), almeno una volta, in tre mesi.

b. Definizione principale: Lettori giorno medio

Quanti leggono o sfogliano il quotidiano X (non importa quale numero), almeno una volta, in media (nel corso di una settimana), in un giorno. Per i quotidiani sportivi vengono utilizzate anche due altre definizioni: Lettori ultimo lunedì e Lettori giorno medio escluso lunedì.

c. Definizioni di frequenza:

Lettori per classi di frequenza: Quanti leggono o sfogliano il quotidiano X, nel corso di tre mesi: 4 o più giorni alla settimana (alta), 1-3 giorni la settimana (media) / meno di 1 giorno la settimana (bassa).

Lettori per frequenza dettagliata: Quanti leggono o sfogliano il quotidiano X: 7 giorni alla settimana / 6 giorni la settimana / 5 giorni la settimana / 4 giorni la settimana / 3 giorni la settimana / 2 giorni la settimana / 1 giorno alla settimana / 2-3 volte in 1 mese / 1 volta al mese / occasionalmente in 3 mesi.

d. Definizioni per singoli giorni della settimana (nei 7 giorni della settimana): Lettori ultima domenica/ultimo sabato/ultimo venerdì/ ecc.

Quanti leggono o sfogliano il quotidiano X (non importa quale numero), almeno una volta, rispettivamente: la domenica / il sabato / il venerdì, ecc.

e. Definizioni secondarie: Lettori ultimi 30 giorni / Lettori ultimi 7 giorni

Quanti leggono o sfogliano il quotidiano X (non importa quale numero), almeno una volta, rispettivamente: in 30 giorni / in 7 giorni.

f. Definizioni collettive: Lettori di quotidiani: nel complesso / ultimi tre mesi / ultimi 30 giorni / ultimi 7 giorni / giorno medio.

Quanti leggono o sfogliano almeno un quotidiano almeno una volta, rispettivamente: in tre mesi / in 30 giorni / in 7 giorni / in un giorno. (L'espressione "quotidiani nel complesso" usata nelle tavole, si riferisce all'estensione a "qualunque quotidiano", iscritto o meno all'indagine. Quando tale espressione manca, la definizione collettiva si intende riferita alle sole testate iscritte.)

Per i PERIODICI E I SUPPLEMENTI DI QUOTIDIANI

a. Definizione-tetto: Lettori totali

Quanti leggono o sfogliano la testata X (non importa quale numero), almeno una volta, in tre mesi se settimanale o supplemento settimanale di quotidiano, oppure in dodici mesi se mensile.

b. Definizione principale: Lettori ultimo periodo

Quanti leggono o sfogliano la testata X (non importa quale numero), almeno una volta, in 7 giorni se settimanale o supplemento settimanale di quotidiano, oppure in 30 giorni se mensile.

c. Definizioni di frequenza:

Lettori per classi di frequenza: Quanti leggono o sfogliano la testata X, nel corso del periodo lungo (tre mesi se settimanale o supplemento settimanale di quotidiano / dodici mesi se mensile), con la seguente frequenza: 9-12 numeri su 12 (alta) / 4-8 numeri su 12 (media) / fino a 3 numeri su 12 (bassa).

Lettori per frequenza dettagliata: Quanti leggono o sfogliano la testata X nel corso del periodo lungo: 12 numeri in 3/12 mesi / 11 numeri in 3/12 mesi / 10 numeri in 3/12 mesi / 9 numeri in 3/12 mesi / 8 numeri in 3/12 mesi / 7 numeri in 3/12 mesi / 6 numeri in 3/12 mesi / 5 numeri in 3/12 mesi / 4 numeri in 3/12 mesi / 3 numeri in 3/12 mesi / 2 numeri in 3/12 mesi / 1 numero in 3/12 mesi o meno.

d. Definizioni collettive: Lettori di settimanali / Lettori di mensili / Lettori di periodici: Lettori totali / Lettori ultimo periodo.

Quanti leggono o sfogliano almeno una delle testate della categoria menzionata (settimanali/mensili/periodici), facendo riferimento alla definizione di Lettori totali e Lettori ultimo periodo.

Per periodici si intendono convenzionalmente testate settimanali e mensili.

8.2.6. Ripartizione ed esecuzione delle interviste

Le interviste vengono ripartite fra le settimane di rilevazione, che possono essere 10 o 11 per ciascun ciclo.

L'equipartizione per settimana opera a livello regionale e, ove possibile, anche a livelli sub-regionali (grandi province, gruppi di province piccole).

L'equipartizione per settimana opera a livello regionale e, ove possibile, anche a livelli sub-regionali (grandi province, gruppi di province piccole).

Sistema CAPI per l'esecuzione e la trasmissione delle interviste

L'intervistatore viene dotato di un PC portatile, nel quale è installato il questionario per le interviste che gli sono state assegnate, nonché i programmi specializzati per la gestione delle interviste e la trasmissione delle stesse – per via telematica – dall'intervistatore al centro di elaborazione dati dell'Istituto. L'intervistatore legge le domande che compaiono man mano sul video del computer e registra le risposte digitando sulla tastiera o mediante il mouse.

L'intervistatore viene dotato anche di uno schermo a colori (tablet), collegato al PC portatile tramite tecnologia Wi-Fi, con lo scopo di visualizzare a favore dell'intervistato i logo delle testate, mentre la scelta delle immagini da mostrare viene guidata da programmi installati sul PC portatile.

Nella gestione dell'intervista, l'intervistatore viene assistito dal PC, specialmente nella somministrazione delle domande di screening, per la selezione delle testate lette nell'ultimo periodo lungo, e per il successivo passaggio – solo per le testate selezionate – alla fase di sviluppo delle domande sulla lettura.

Sistema CAWI per l'esecuzione e la trasmissione delle interviste

A ciascun individuo selezionato per la rilevazione online viene spedita un'e-mail di invito a partecipare all'indagine, nella quale è presente il link per collegarsi al questionario. Il software di gestione dell'intervista (sistema CAWI) guida l'intervistato nelle prime domande, volte a verificare l'eleggibilità al target d'intervista. Se l'individuo supera questa fase preliminare, inizia il questionario vero e proprio, basato sulle stesse domande utilizzate nell'intervista face to face.

Il questionario in versione CAWI ricalca tutte le logiche e i requisiti di coerenza presenti nel questionario in versione CAPI, così come l'esatta specifica dei criteri di controllo, ma è implementato ad hoc per l'autosomministrazione, con un layout grafico piacevole e intuitivo, che rende chiare e comprensibili le diverse domande e le possibili risposte e semplice la compilazione.

Il flusso del questionario e la registrazione assistita e controllata delle risposte sono gestiti in automatico dal software di intervista.

Addestramento degli intervistatori

L'esecuzione delle interviste personali face to face CAPI viene preceduta da riunioni di formazione degli intervistatori (briefing), condotte via webinar a cura dell'Istituto. Vengono inoltre realizzati briefing telefonici di ripasso, studiati ad hoc per i singoli intervistatori.

Il programma del briefing è il seguente:

- illustrazione delle procedure di estrazione dei nominativi dalle liste elettorali e gestione della lista dei nominativi per le interviste e per le eventuali sostituzioni.

- illustrazione del questionario e delle tecniche di intervista.
- esemplificazione mediante una intervista simulata, durante la riunione.
- effettuazione di una intervista di prova. Ciascun intervistatore viene chiamato ad effettuare parte dell'intervista simulata, al fine di verificare l'apprendimento dell'utilizzo di PC e Tablet.
- discussione con i singoli intervistatori sull'intervista di prova fatta e sulle eventuali difficoltà incontrate.

Le riunioni di briefing, che vengono tenute da istruttori che sono scelti dall'Istituto tra i propri ricercatori addetti all'Indagine "Audicom Print" o funzionari del reparto field, hanno la durata di almeno mezza giornata.

Controllo del lavoro degli intervistatori

Un controllo del lavoro di rilevazione di ciascun Istituto viene effettuato, con i propri criteri, da un Istituto di certificazione esterno, nel quadro dei controlli che questo Istituto effettua su tutte le fasi di lavoro (campionamento, rilevamento ed elaborazione).

Indipendentemente dai controlli suddetti e comunque in via preliminare a quelli che riguardano le interviste, l'istituto esegue i seguenti controlli.

I controlli sulla qualità del lavoro svolto dagli intervistatori vengono effettuati a due livelli:

1. A tavolino, esaminando criticamente il materiale di rilevazione inviato dall'intervistatore (questionari, ecc.).
2. Sul campo, effettuando interviste di controllo telefoniche alle persone che risultano intervistate.

I controlli a tavolino hanno per oggetto:

- i questionari compilati, per controllare la correttezza formale delle registrazioni fatte, rilevare l'eventuale presenza di errori, incoerenze o lacune di compilazione. Gli errori o omissioni che fossero rimasti dopo le correzioni effettuate durante l'intervista a seguito dell'intervento del programma CAPI (cleaning in tempo reale) vengono individuati e corretti mediante ulteriori programmi di cleaning ex-post. Viene rilevata anche la presenza di annotazioni speciali su anomalie o incidenti nella conduzione delle interviste o nelle condizioni ambientali dell'intervista, che l'intervistatore è tenuto a fare in margine al questionario;
- i moduli contenenti i nominativi estratti dalle liste elettorali ed utilizzati (o non utilizzati) per fare le interviste, per controllare la correttezza dell'utilizzazione in relazione alle varie siglature concernenti la gerarchia di utilizzo (Ordinari e di Riserva);
- i documenti dell'operazione di estrazione, per controllare (anche attraverso confronti tra le informazioni rilevate all'ufficio elettorale, le istruzioni impartite e le informazioni presenti sui questionari e gli altri documenti dell'intervista) la correttezza delle operazioni svolte presso l'ufficio elettorale, ed in particolare:

- la coincidenza dei numeri di sezione e di nominativo assegnati con quelli delle sezioni e dei nominativi effettivamente estratti;
- il rispetto delle regole per l'estrazione e per le sostituzioni delle sezioni o di nominativi non estraibili;
- la completezza delle registrazioni e verbalizzazioni fatte, anche dopo l'intervista, sul modulo di estrazione.

In caso di errore o di sospetto di errore, si indaga ulteriormente anche sulla buona fede dell'intervistatore nel commetterlo, facendo anche telefonate di controllo all'ufficio elettorale (p. es. sull'effettiva presenza nel registro elettorale di un dato nominativo che figura sia estratto che intervistato). L'indagine viene approfondita fino alla rilevazione di eventuali cambiamenti avvenuti nel periodo trascorso tra la rilevazione del nominativo e l'intervista.

I controlli sul campo mediante intervista telefonica vengono fatti da intervistatori specializzati (che vengono adeguatamente coinvolti nella problematica dei controlli e ne conoscono un'ampia casistica).

Il controllore chiama il numero telefonico registrato dall'intervistatore sul questionario come quello della famiglia in cui è stata fatta l'intervista e prende contatto con la persona che risulta intervistata. Il numero telefonico, se non è stato indicato dall'intervistatore, viene ricercato in base al nominativo e all'indirizzo.

All'intervistato vengono rivolte domande per accertare quanto segue:

- se ricorda di essere stato visitato da un intervistatore/trice "che Le ha rivolto domande sulle abitudini di lettura di giornali (quotidiani, supplementi di quotidiani, settimanali, mensili)". (Se no, vengono fatti tutti gli approfondimenti del caso).
- se l'intervistatore era un uomo o una donna
- se l'intervista ha avuto luogo nell'abitazione o in altro luogo e quale (in particolare: se si trattava di un luogo privato o di pubblico passaggio; nel caso di luoghi di lavoro: se di proprietà dell'intervistato o della famiglia)
- se l'intervistatore aveva un tablet sulla quale comparivano i "logo di giornali o riviste"
- se durante l'intervista era presente qualche altra persona e, in caso affermativo, chi, e se alle domande ha risposto sempre l'intervistato personalmente o anche l'altra persona (in questo secondo caso vengono fatti approfondimenti per accertare se l'intervento dell'altra persona abbia riguardato anche le risposte da dare alle domande sulla lettura: in questo caso l'intervistatore avrebbe consentito un intervento che doveva essere evitato)
- oltre al sesso, all'età, al grado di istruzione e alla condizione professionale dell'intervistato, si rileva il numero, il sesso, l'età e la condizione professionale dei componenti la famiglia.

L'accertamento di errori, anomalie o interrogativi risultanti dai controlli a tavolino o dalle interviste telefoniche di controllo, dà origine a telefonate all'intervistatore per chiedere spiegazioni, ove necessario, e per far notare errori o imperfezioni e raccomandarne la non ripetizione per l'avvenire.

In caso di errore, la singola intervista viene annullata ed esclusa dalle elaborazioni dei dati.

Nei casi gravi (ripetizione di errori in più interviste), viene annullato il lavoro dell'intero incarico dell'intervistatore.

I controlli descritti sono estesi a tutti gli intervistatori e a tutte le interviste.

Le interviste di controllo telefoniche sono estese a tutti gli intervistatori e riguardano tutti i CDC eseguiti e tutte le interviste di ciascun CDC.

Controllo delle interviste online

I controlli sulla qualità delle interviste online vengono svolti a tavolino, esaminando criticamente i questionari compilati per controllare la correttezza formale delle registrazioni fatte, rilevare l'eventuale presenza di errori e incoerenze di compilazione. Gli errori o omissioni che fossero rimasti dopo le correzioni effettuate durante l'intervista a seguito dell'intervento del programma CAWI (cleaning in tempo reale) vengono individuati mediante ulteriori programmi di cleaning ex-post.

Un ulteriore controllo viene effettuato sulle principali informazioni di profilo dell'intervistato (sesso, età, territorio, numero di componenti della famiglia): viene controllata la coerenza fra quanto dichiarato in intervista e quanto presente nell'anagrafica compilata dall'individuo in fase di registrazione all'access panel.

8.3. Elaborazioni

Fasi di elaborazione per la produzione dei risultati

Le fasi di elaborazione sono finalizzate alla creazione dei seguenti prodotti:

- i rapporti con le tavole statistiche contenenti le stime dell'indagine, quale base di conoscenza e documentazione dei risultati e quale mezzo principale di diffusione dei dati (disponibili sui siti www.audipress.it e www.audicom.net);
- l'archivio su supporto informatico detto "nastro di pianificazione", con i dati in forma adatta per essere riprodotti, selezionati e rielaborati, attraverso sistemi computerizzati, secondo le esigenze dell'utente.

Il nastro di pianificazione Print viene utilizzato anche come base dati funzionale alla successiva fase di fusione con il dato digital.

L'elaborazione dell'indagine "Audicom Print" prevede la cumulazione di tre cicli di rilevazione single source.

Le fasi di controllo e cleaning dei dati raccolti attraverso i questionari vengono svolte, a cura dell'Istituto esecutore, entro il quadro dell'attività di rilevamento e messa a punto dei dati originali.

Le fasi principali di trattamento e produzione dei dati sono, nell'ordine le seguenti:

- a) l'applicazione dell'algoritmo di correzione del BIAS della piattaforma di erogazione dell'intervista
- b) la "ponderazione"
- c) il trattamento delle mancate risposte alle variabili Reddito e Classe Socio-Economica nei dati rilevati online
- d) la procedura di armonizzazione CAWI - CAPI delle variabili Reddito e Classe Socio-Economica
- e) il trattamento delle mancate uscite dei quotidiani
- f) l'attribuzione dei "contatti"
- g) i "pattern di lettura"

○*○*○*○*○*○*○

a) L'applicazione dell'algoritmo di correzione del bias della piattaforma di erogazione dell'intervista

In fase di analisi ed elaborazione dei risultati rilevati con tecnica CAWI viene applicata in via preliminare una procedura volta a riconciliare i dati ottenuti con la piattaforma di erogazione CAWI con quelli ottenuti via CAPI (che consideriamo «source of truth»).

L'algoritmo di correzione del bias della piattaforma di erogazione dell'intervista permette di calcolare la probabilità di compliance dei questionari rilevati online, determinata in base al numero di testate lette (in versione cartacea o digitale replica) e alla durata di intervista.

Si basa inoltre sui seguenti criteri operativi:

- un algoritmo che opera a livello "micro" e non "macro", respondent-level (non solo sui risultati aggregati)
- variabili discrete ri-stimate dalla procedura sull'indagine CAWI, distribuite sugli stessi livelli utilizzati nelle rilevazioni CAPI e CAWI
- criterio oggettivo di qualità dell'intervista, identificato nella durata dell'intervista legata al numero di testate/versioni lette

Sulla base di tali premesse questo algoritmo è applicato alle interviste CAWI dei tre cicli di rilevazione considerati; si tratta di un algoritmo esponenziale che calcola la probabilità di compliance di ciascuna intervista, che può assumere un valore compreso tra 0 e 1.

Da un punto di vista procedurale, i dati vengono trattati separatamente per Quotidiani e Periodici. Da un punto di vista strutturale, l'algoritmo viene applicato a tutti i target rilevati online: al target di individui maggiorenni, al target Giovani 13-17 anni e al target CSS, applicando un adattamento che tiene conto della natura di questo segmento di popolazione.

Le variabili utilizzate dall'algoritmo sono le seguenti:

- il numero totale di Quotidiani e Supplementi letti nel periodo allargato in versione Carta
- il numero totale di Quotidiani e Supplementi letti nel periodo allargato in versione Replica
- il numero totale di Settimanali letti nel periodo allargato in versione Carta
- il numero totale di Settimanali letti nel periodo allargato in versione Replica
- il numero totale di Mensili letti nel periodo allargato in versione Carta
- il numero totale di Mensili letti nel periodo allargato in versione Replica
- il tempo in minuti impiegato per portare a termine il questionario.

Al termine dell'applicazione dell'algoritmo, l'insieme delle interviste CAWI viene suddiviso in tre gruppi:

- interviste considerate compliant, che vengono mantenute nel campione con peso iniziale pari all'unità
- interviste considerate potenzialmente non compliant, che vengono mantenute nel campione ma con un peso iniziale, a monte della ponderazione, pari alla probabilità di compliance, calcolata dall'algoritmo
- interviste che vengono considerate ragionevolmente non compliant, cui viene attribuito un peso pari a zero.

A valle di questa operazione, il campione complessivo (CAWI e CAPI) viene riponderato sulle caratteristiche demografiche, applicando la procedura di ponderazione (vedi par. 4.3), in modo tale da ripristinare la rappresentatività del campione rispetto alla popolazione di riferimento.

b) La ponderazione

La ponderazione viene eseguita allo scopo di riportare nel campione le proporzioni dell'universo di riferimento (popolazione adulta italiana di 13 anni e oltre secondo vari caratteri).

La ponderazione è necessaria in **via primaria** perché bisogna: 1) ripristinare con precisione la proporzionalità geografica all'universo, che viene modificata volutamente a priori con le operazioni di sovracampionamento, che riguardano alcune province; 2) ripristinare con precisione la proporzionalità dell'universo secondo il carattere "CSS/non CSS" all'interno delle singole regioni o areole, perché questo segmento della popolazione è stato fatto oggetto di sovracampionamento; 3) ripristinare con precisione la proporzionalità dell'universo secondo il carattere titolo di studio, perché naturalmente più elevato nella quota di interviste rilevate online.

La ponderazione offre inoltre, in **via secondaria**, l'occasione di apportare alcuni miglioramenti alla rappresentatività del campione, in quanto consente di rimediare ad eventuali anomalie o imperfezioni nella composizione del campione per i seguenti caratteri (noti per l'universo) degli individui: il sesso, l'età, il ruolo di responsabile acquisti e quello di capofamiglia.

I caratteri geografici tenuti sotto controllo sono:

- le singole areole (71)
- comune capoluogo e non capoluoghi, entro l'areola (areola = 71 singole province o gruppi di province della stessa regione)
- classi di ampiezza demografica dei comuni, entro le 18 regioni

In ogni areola vengono tenuti sotto controllo i seguenti caratteri individuali:

- sesso incrociato con classi di età
- ruolo di capofamiglia e non
- ruolo di responsabile acquisti e non
- caratteristica di CSS e non
- ciclo di rilevazione

In ogni area geografica viene tenuto sotto controllo il carattere titolo di studio.

Al termine della ponderazione si dispone, per ogni singolo individuo, di un fattore di espansione che, applicato alle informazioni provenienti dall'individuo, consente di produrre stime in valori riferibili all'universo.

Tutte le stime vengono riportate ad un unico universo di individui adulti, dai 13 anni in su, che viene espresso in migliaia.

c) Trattamento delle mancate risposte al reddito e alla classe socio-economica nei dati rilevati online

Nelle interviste autocompilate online, per le informazioni familiari relative al reddito netto mensile complessivo e alla classe socioeconomica viene lasciata all'intervistato la facoltà di non rispondere.

Nell'intervista personale face to face CAPI, tali informazioni sono stimate dall'intervistatore e pertanto sempre presenti in forma completa.

Le informazioni mancanti nelle interviste online devono essere ricostruite in fase di elaborazione dei dati.

A tal fine viene utilizzata una procedura di merging che permette di trovare coppie di individui in modo controllato, cioè tra loro somiglianti per caratteristiche sociodemografiche, geografiche e di composizione della famiglia.

Interviste in cui operare la ricostruzione

Le interviste in cui è necessario intervenire con la ricostruzione dell'informazione sul Reddito o sulla Classe Socio-Economica sono quelle rilevate online in cui l'informazione risulta mancante (l'intervistato ha risposto "non indico") in almeno una delle due domande.

In particolare, la ricostruzione della risposta mancante avviene attribuendo l'informazione a partire da un individuo "donatore" presente nel campione intervistato personalmente, che sia gemello dell'individuo "ricevente" intervistato online, secondo il seguente criterio.

- Se nell'intervista rilevata online manca solo l'informazione del reddito, la procedura di merging attribuisce solo questa informazione, lasciando l'informazione originale sulla classe socioeconomica;
- Se nell'intervista manca solo l'informazione della classe socioeconomica la procedura di merging attribuisce solo questa informazione, lasciando l'informazione originale sul reddito;
- Se nell'intervista mancano sia l'informazione del reddito sia della classe socioeconomica la procedura di merging attribuisce entrambe le informazioni.
-

d) Procedura di armonizzazione CAWI-CAPI delle variabili reddito e classe socio-economica

Al fine di armonizzare il profilo di reddito e della classe socioeconomica rilevato nei due campioni CAWI e CAPI, viene utilizzata una procedura di elaborazione, che permette:

- di allineare nella CAWI questi due caratteri rispetto a quanto stimato dall'intervistatore nella rilevazione CAPI
- di preservare al contempo la natura della distribuzione originale di queste variabili nella CAWI

Nel dettaglio, l'elaborazione prevede l'applicazione di due distinte procedure di merging per soggetto, che utilizzano come forti variabili di aggancio tra i due campioni CAWI e CAPI le

autovalutazioni del reddito e della classe socioeconomica. In aggiunta, le procedure utilizzano altre variabili d'aggancio:

- individuali, come il titolo di studio, il sesso, l'età e l'area geografica
- familiari, come il numero di componenti che lavorano, il numero di componenti, la presenza di almeno un componente appartenente alla categoria CSS.

Per ciascuna delle due variabili (reddito e classe socioeconomica), dopo aver identificato gli individui «gemelli» tra CAWI e CAPI, viene applicato al gemello CAWI il differenziale tra l'autodichiarazione che l'individuo gemello CAPI ha dato della variabile e la stima fatta dall'intervistatore.

Questa procedura viene applicata solamente nei casi in cui il differenziale assume valore diverso da zero, ovvero nei casi in cui nella CAPI non ci sia coincidenza tra autodichiarazione dell'individuo e stima fatta dall'intervistatore.

e) Trattamento delle mancate uscite di quotidiani

Se durante il periodo di rilevazione in alcuni giorni un quotidiano non viene pubblicato, viene apportata ai lettori giorno medio "carta", ai lettori giorno medio "carta e/o replica" e ai lettori giorno medio "replica" una correzione sistematica dell'effetto che le mancate uscite potrebbero avere avuto sulla stima dei "lettori giorno medio". Questa correzione viene apportata con il metodo seguente: quando la lettura di una testata ha luogo in un giorno che è stato segnalato come di mancata uscita della testata, tale lettura non viene considerata ai fini delle stime per nessuna delle definizioni di lettura. In particolare, per il calcolo dei "lettori nel giorno medio" il giorno di mancata uscita viene sottratto sia al numeratore che al denominatore della frazione su cui è basato tale calcolo.

Se durante il periodo di rilevazione in alcuni giorni un quotidiano non viene pubblicato, viene apportata ai lettori giorno medio "carta", ai lettori giorno medio "carta e/o replica" e ai lettori giorno medio "replica" una correzione sistematica dell'effetto che le mancate uscite potrebbero avere avuto sulla stima dei "lettori giorno medio". Questa correzione viene apportata con il metodo seguente: quando la lettura di una testata ha luogo in un giorno segnalato come di mancata uscita della testata, tale lettura non viene considerata ai fini delle stime per nessuna delle definizioni di lettura. In particolare, per il calcolo dei "lettori nel giorno medio" il giorno di mancata uscita viene sottratto sia al numeratore che al denominatore della frazione su cui è basato tale calcolo.

Ricordiamo che il rapporto utilizzato per il calcolo dei "lettori nel giorno medio", nel caso di quotidiani che escono 7 giorni su 7, considerato il periodo di 7 giorni precedenti il giorno di intervista, è il seguente:

lettura nel 1° giorno + lettura nel 2° giorno + ...(ecc.) + lettura nel 7° giorno

Nel caso che al 2° giorno corrisponda una mancata uscita, la frazione viene modificata nel modo seguente: al numeratore si toglie la "lettura nel 2° giorno" (cioè si toglie un addendo su 7) e si riduce di uno il denominatore (che diventa 6).

Questa formula consente di produrre, per quanto riguarda la definizione "lettori giorno medio", stime più vicine ai valori che si sarebbero probabilmente ottenuti "se le mancate uscite non si fossero verificate".

f) Attribuzione dei "contatti"

Quotidiani - Calcolo ed estensione dell'informazione contatti medi per copia media

CALCOLO

L'informazione unica, raccolta in sede di intervista, "Pensi all'ultimo giorno ... in cui ha letto o sfogliato ... Quante volte ... ha preso in mano quella copia del quotidiano ..." viene trattata per generare un'unica informazione: numero di contatti medi per copia media.

L'unico dato disponibile genera automaticamente, salvo una valorizzazione dell'ultima classe aperta (5 o più volte), l'informazione: numero di contatti medi per copia media.

I pesi attribuiti a ciascuno dei punti della scala sono i seguenti:

"Quante volte...":

- 1 volta	peso 1.00
- 2 volte	peso 2.00
- 3 volte	peso 3.00
- 4 volte	peso 4.00
- 5 o più volte	peso 5.50

Il peso così attribuito rappresenta l'informazione desiderata.

ESTENSIONE

Il questionario dell'indagine "Audicom Print" raccoglie l'informazione sopra citata esclusivamente per i lettori ultimi 7 giorni, rendendo così possibile il calcolo dell'informazione sintetica "contatti medi per copia media" solo per tali lettori. Tuttavia, poiché l'informazione "contatti medi per copia media" deve essere disponibile per gli usi operativi di pianificazione, è indispensabile che essa sia estesa a tutti i lettori del periodo lungo. Per questo motivo viene eseguita la cosiddetta estensione dei contatti medi per copia media.

L'estensione dei valori viene effettuata a partire dai valori calcolati per i lettori ultimi 7 giorni che hanno dichiarato almeno 1 giorno di lettura.

Il valore viene esteso all'interno delle celle di popolazione definite sia dai caratteri di lettura (frequenza di lettura) e sia dai caratteri geo-sociodemografici, attribuendo ad ogni individuo lettore periodo lungo, appartenente ad una determinata cella, il valore medio dell'informazione calcolata, all'interno della stessa cella, per i lettori ultimi 7 giorni che hanno dichiarato almeno 1 giorno di lettura.

Nel caso in cui una cella contenga esclusivamente lettori del periodo lungo che non sono lettori ultimi 7 giorni con almeno 1 giorno di lettura dichiarato (così che venga a mancare il valore calcolato da dichiarazione), si opera passando al livello superiore della segmentazione in celle.

Al termine del processo tutti i lettori del periodo lungo (quindi anche coloro che non sono lettori ultimi 7 giorni...) dispongono dell'informazione sintetica contatti medi per copia.

Periodici - Calcolo ed estensione dell'informazione contatti medi per pagina media

CALCOLO

Le due informazioni, raccolte distintamente in sede di intervista, "in questi ultimi 7/30 giorni quante volte le è capitato di leggere o sfogliare ..." e "Pensi all'ultima volta ... Quante pagine ha letto o sfogliato ..." vengono trattate mediante ponderazione per generare un'unica informazione: numero di contatti medi per pagina media.

I pesi attribuiti a ciascuno dei punti delle scale delle due domande sono i seguenti.

"Quante volte...":

- 1 volta	peso 1.00
- 2 volte	peso 2.00
- 3 volte	peso 3.00
- 4 volte	peso 4.00
- 5 volte o più	peso 7.75

"Quante pagine...":

- tutte le pagine o quasi	peso 0.975
- circa 3/4	peso 0.750
- circa la metà	peso 0.500
- circa 1/4	peso 0.250
- meno di 1/4	peso 0.075

I pesi così definiti sono quindi moltiplicati tra loro in modo da ottenere l'informazione sintetica desiderata.

ESTENSIONE

Nel questionario dell'indagine "Audicom Print", le domande sopra citate sono poste esclusivamente ai lettori ultimo periodo, i quali rendono così possibile il calcolo dell'informazione sintetica "contatti medi per pagina media". Tuttavia, tenuto conto che l'informazione "contatti medi per pagina media" deve essere disponibile per gli usi operativi di pianificazione, è indispensabile che essa sia estesa a tutti i lettori del periodo lungo.

Per questo motivo viene eseguita la cosiddetta estensione dei contatti medi per pagina, partendo dai valori calcolati per gli individui lettori ultimo periodo.

L'estensione dei valori avviene all'interno di celle di popolazione definite sia da caratteri di lettura (frequenza di lettura) e sia da caratteri sociodemografici, attribuendo ad ogni individuo lettore periodo lungo, appartenente ad una determinata cella, il valore medio dell'informazione calcolata all'interno della stessa cella per i lettori ultimo periodo.

Nel caso in cui una cella contenga esclusivamente lettori del periodo lungo che non sono lettori ultimo periodo (così che venga a mancare il valore calcolato da dichiarazione) si opera passando al livello superiore individuato dalla segmentazione in celle.

Al termine del processo tutti i lettori del periodo lungo (quindi anche coloro che non sono lettori ultimo periodo) dispongono dell'informazione contatti medi per pagina media. Tale informazione viene riportata per ogni individuo lettore del periodo lungo di una certa testata nel nastro di pianificazione e resa disponibile ai fini di analisi e di pianificazione del mezzo stampa.

g) "Pattern di lettura"

Premessa

I pattern di lettura sono prodotti per tutte le testate rilevate e pubblicate nell'indagine per essere aggiunti alle altre informazioni già presenti all'interno del "nastro di pianificazione".

Per pattern di lettura si intende una sequenza di valori logici ($S_i=1$; $No=0$) associati alla possibile esposizione a differenti uscite di una testata da parte dello stesso individuo lettore.

Ad esempio, la sequenza 101000100101 va interpretata per un determinato individuo ed una certa testata come:

- a. 5 uscite viste su 12 complessive, la 1a, 3a, 7a, 10a e 12a
- b. 7 uscite non viste, la 2a, 4a, 5a, 6a, 8a, 9a e 11a²

Il calcolo della copertura netta e della distribuzione di frequenza ne risulta assai semplificato, poiché si riduce ad un semplice conteggio dei valori ($1=S_i$) all'interno del pattern. Così delle prime 4 uscite verranno viste dal lettore in esame solo 2, la prima e la terza.

Come si vede, oltre a semplificare le operazioni, l'uso dei pattern introduce una visione posizionale della esposizione cumulata nel tempo.

Essi forniscono come informazione non solo quante uscite verranno viste, ma anche esattamente in quali momenti. *(Ciò vale non con riferimento alla lettura di specifici numeri di una testata, ma con riferimento alla possibile lettura di un qualsiasi numero, anche arretrato, nelle varie unità di tempo considerate (sequenza di 28 giorni per i quotidiani, di 12 settimane per i settimanali e i supplementi settimanali di quotidiani, di 12 mesi per i mensili)).*

I vantaggi dei pattern di lettura

Come già detto, nel “nastro di pianificazione” i pattern vengono aggiunti a tutte le informazioni da sempre fornite al mercato.

Questi, in sintesi, i dati di lettura presenti nel “nastro di pianificazione” per ogni individuo/testata:

	<i>Quotidiani</i>	<i>Periodici</i>
<i>Lettura allargata</i>	Ultimi 3 mesi	Ultimi 3 mesi/12 mesi
<i>Lettura media</i>	Giorno medio	Ultimo periodo
<i>Frequenza di lettura</i>	Da 1 giorno su 7 a 7 giorni su 7 Altri lettori in 7 giorni Lettori 1 volta al mese Lettori 2-3 volte in 3 mesi	Da 1 numero su 12 a 12 numeri su 12
<i>Pattern di lettura</i>	28 uscite (4 settimane)	12 uscite

Rendendo trasparente la sequenza temporale dell'esposizione, i pattern si prestano facilmente ad usi sofisticati quali lo studio della memorizzazione di una comunicazione e dell'efficacia di una campagna tramite appositi algoritmi (ad esempio Morgenstern).

La costruzione dei pattern di lettura

PRINCIPI GENERALI

- Per consentire i calcoli di pianificazione basati sui “pattern di lettura”, garantendo al tempo stesso la coerenza con le altre variabili normalmente utilizzate nei processi di pianificazione, viene costruita anzitutto una nuova variabile strumentale detta *frequenza di lettura normalizzata*, che tiene conto sia della frequenza di lettura dichiarata, sia del valore di probabilità di lettura, sia degli obiettivi di rappresentazione o simulazione voluti con i pattern. Tale variabile corrisponde, in misura precisa e coerente con la probabilità matematica di lettura ma in forma *discreta* (o *discontinua*, cioè per valori interi e non

continui), al *numero delle uscite* a cui l'individuo lettore del periodo lungo di una determinata testata è esposto.

- In coerenza con la frequenza di lettura normalizzata, i “pattern di lettura” corrispondono al posizionamento delle uscite a cui i vari lettori sono esposti. Questo posizionamento viene generato scegliendo casualmente tra tutte le numerosissime sequenze possibili e facendo in modo che tutte le possibili sequenze vengano considerate come ugualmente probabili, senza privilegiare alcuna particolare uscita, tenuto conto che la dichiarazione di frequenza dichiarata è comunque generica e cioè non presuppone alcuna combinazione particolare di uscite nel periodo lungo.
- Condizione essenziale della produzione dei “pattern di lettura” è che venga riprodotto, con riferimento agli esposti attribuiti ad ogni uscita, il giusto profilo geo-sociodemografico dei lettori ultimo periodo, considerando l'intero bacino del periodo lungo e controllando al tempo stesso la coerenza con il numero totale di uscite cui è esposto potenzialmente il lettore.

Per i QUOTIDIANI

- Il pattern copre 4 settimane (28 giorni/uscite) corrispondenti ad 1 mese convenzionale. I lettori con frequenza di lettura superiore ai 30 giorni sono attribuiti ad una uscita convenzionale che rappresenta il 29mo punto del pattern.
- Ogni testata letta nel periodo lungo ha almeno un punto (uno qualsiasi dei 29 possibili) del pattern posto uguale a 1. I lettori di almeno una uscita su 28 calcolati da pattern corrispondono ai lettori potenziali dichiarati negli ultimi 30 giorni, mentre il totale dei lettori esposti ad almeno 1 uscita su 29 (compresa l'ultima convenzionale) corrisponde ai lettori del periodo lungo.
- Per ogni testata, a livello di uomini, donne e responsabili acquisti, le stime di lettura ponderate calcolate per ciascuna uscita tra quelle possibili devono essere statisticamente equivalenti a quelle del giorno medio. Nel processo che realizza tale condizione si tiene anche conto della necessità di ottimizzare, uscita per uscita, la rappresentazione geo-sociodemografica dei lettori di ogni uscita in rapporto ai lettori del giorno medio.
- Una fase di controllo finale verifica il raggiungimento degli obiettivi, compreso il fatto che la media delle letture ottenute con i 28 punti del pattern deve corrispondere alla lettura giorno medio.

Per i PERIODICI

- Il pattern copre 12 settimane per i settimanali e i supplementi settimanali di quotidiani, 12 mesi per i mensili, tante quante sono le classi delle dichiarazioni di frequenza del questionario.

- Ogni testata letta nel periodo lungo ha almeno un punto (uno qualsiasi dei 12 possibili) del pattern posto uguale a 1, cioè i lettori di almeno 1 numero su 12 calcolati dal pattern corrispondono ai lettori potenziali dichiarati negli ultimi 3/12 mesi.
- Ad ogni testata letta per un certo numero di uscite corrisponde una scelta casuale di uscite tra quelle possibili, cui l'individuo è esposto. La somma di tali uscite restituisce la frequenza normalizzata dell'individuo.
- Per ogni testata, a livello di uomini, donne e responsabili acquisti, le stime di lettura ponderate calcolate per ciascuna uscita tra quelle possibili devono essere statisticamente equivalenti a quelle dell'ultimo periodo. Nel processo che realizza tale condizione si tiene anche conto della necessità di ottimizzare, uscita per uscita, la rappresentazione geo-sociodemografica dei lettori di ogni uscita in rapporto ai lettori dell'ultimo periodo.
- Una fase di controllo finale verifica il raggiungimento degli obiettivi, compreso il fatto che la media delle letture ottenute con i 12 punti del pattern deve corrispondere alla lettura ultimo periodo.

8.3.1. Attività di controllo a cura di terze parti

I processi di misurazione della readership sono sottoposti periodicamente a procedure di verifica ad opera di terze parti indipendenti (Reply Consulting S.r.l.), nell'ambito dell'incarico conferito da Audicom.

8.4. Definizione e note per l'uso dei dati

Ecco alcune precisazioni sui caratteri usati per l'analisi statistica dei lettori, sia per il target "Lettori carta e/o replica" che per il target "Lettori carta". Per l'analisi statistica del target "Lettori replica" viene utilizzato un sottoinsieme di questi stessi caratteri.

Soglie minime di pubblicazione dei dati

La pubblicazione del dato di ciascuna testata è subordinata al raggiungimento delle seguenti soglie minime:

- 75.000 (settantacinquemila) lettori "ultimi 7 giorni" per i quotidiani (testate regolarmente registrate comprese le eventuali edizioni locali);
- 90.000 (novantamila) lettori "ultimo periodo" per i periodici.

Le suddette soglie minime vanno considerate indipendentemente dagli intervalli fiduciari. I dati sottosoglia pertanto non potranno essere pubblicati e l'Editore dei mezzi i cui dati fossero sottosoglia non potrà utilizzarli o parlarli a terzi.

I caratteri di analisi

Responsabile acquisti

All'interno di ogni famiglia viene attribuito il ruolo di "responsabile degli acquisti" (dei beni di consumo corrente) ad una ed una sola persona, che può essere o no l'intervistato medesimo.

Tale definizione differisce in parte dal criterio "donna di casa", nel senso che responsabile degli acquisti può anche essere l'uomo, come di fatto avviene nelle famiglie in cui non vi sono donne, oppure vi sono donne che non si occupano degli acquisti di prodotti alimentari e di prodotti per la casa.

Capofamiglia

All'interno di ogni famiglia viene attribuito, solo ad uno dei componenti della famiglia, che può essere o no l'intervistato, la qualifica di capofamiglia. Il ruolo di responsabile degli acquisti e quello di capofamiglia non sono incompatibili tra loro.

Naturalmente il capofamiglia può essere sia un uomo che una donna.

Classe di età

Ai fini della collocazione nella classe di età si considerano gli anni compiuti alla data dell'ultimo compleanno. La dichiarazione dell'intervistato è controllata con l'anno di nascita riportato sui registri elettorali.

Titolo di studio

Il grado di istruzione è misurato dal titolo di studio effettivamente conseguito dall'intervistato. Di eventuali corsi o anni di studio di grado superiore non si tiene conto se non hanno dato luogo a un nuovo titolo.

Classe di reddito

Il reddito mensile della famiglia viene stimato dall'intervistatore. La valutazione fatta dall'intervistatore del reddito mensile delle famiglie del campione è molto vicina al reddito "speso" dalle famiglie.

Infatti, è stato dimostrato che il valore medio del reddito stimato dagli intervistatori, estrapolato all'universo di tutte le famiglie, porta ad un valore molto prossimo a quello dei "consumi finali interni delle famiglie" nel "conto economico nazionale delle risorse e degli impieghi". Questo significa che l'intervistatore, invitato a valutare la classe di reddito, tende a stimare con buona attendibilità la parte del reddito destinata ai consumi, al netto cioè dei risparmi pluriennali, delle tasse e dei contributi sociali e previdenziali. Da notare che si tratta proprio della parte di reddito più importante per l'uso a cui l'informazione è destinata: classificare i dati secondo la variabile "potere di acquisto immediato della famiglia per l'acquisizione o il godimento di beni di consumo durevole o non durevole, e servizi disponibili sul mercato".

Presenza di bambini e ragazzi

La presenza di bambini o ragazzi fa riferimento alla famiglia cui appartiene l'individuo a cui si riferiscono le letture oggetto dei dati. In altre parole, il dato "lettori della testata X", nella colonna "famiglie con bambini o ragazzi della classe di età Y1-Y2" significa: "quanti lettori della testata X vivono in famiglie in cui sono presenti bambini o ragazzi della classe di età Y1-Y2", indipendentemente dal rapporto di parentela tra i lettori ed i bambini o ragazzi.

Capoluogo e non capoluogo

Nelle tavole in cui vi sono solo le voci "capoluogo" e "non capoluogo" si intende: "capoluogo di regione o di provincia".

Ampiezza demografica dei comuni

I gruppi di analisi considerati nelle elaborazioni fanno riferimento alla classe di popolazione totale (abitanti di tutte le età dei comuni di residenza degli intervistati).

Categoria economico-sociale

L'attribuzione a quella che impropriamente viene chiamata classe sociale è fatta dallo stesso intervistatore. È un indicatore sintetico di stile di vita e al tempo stesso di capacità di spesa. Empiricamente si è verificato che nel valutare la classe sociale l'intervistatore si basa su queste tre variabili: professione del capofamiglia, livello di istruzione, abitazione. Certamente entrano in gioco anche i tradizionali segni di "status", come l'abbigliamento dell'intervistato, il modo di parlare, la facilità di rapporto con l'intervistatore e simili, benché non sia possibile valutare in modo preciso la loro influenza.

Professioni e categorie non professionali

La classificazione analitica delle professioni comprende i seguenti gruppi:

- (1) Imprenditori (datori di lavoro): coloro che gestiscono in conto proprio un'impresa, nella quale non impiegano l'opera manuale propria, ma solo la propria opera direttiva.
- (2) Liberi professionisti: coloro che esercitano in conto proprio una professione con o senza l'aiuto di persone retribuite (avvocati, medici, consulenti fiscali, ecc.).
- (3) Dirigenti: coloro che esercitano, contro retribuzione, una funzione direttiva.
- (4) Quadri o funzionari: coloro che esercitano la propria attività come dipendente di un'azienda (privata o ente pubblico) ad un livello intermedio (non direttivo).
- (5) Ufficiali superiori di un corpo delle forze armate o funzionario di un corpo paramilitare: coloro che hanno ottenuto i gradi dal maggiore in su nell'esercito, nella marina, nell'aviazione, nella polizia, nei carabinieri, nella guardia di finanza, nei vigili del fuoco, nei guardacoste, nel corpo forestale.
- (6) Possidenti: coloro che godono di redditi elevati e per i quali la fonte principale è di natura patrimoniale.
- (7) Impiegati: coloro che esercitano, contro retribuzione, una funzione esecutiva (di concetto o d'ordine), con lavoro più intellettuale che manuale. Le persone che svolgono funzioni assimilabili in parte a quelle impiegatizie e in parte a quelle operaie, le categorie intermedie, sono classificate fra gli impiegati.
- (8) Negozianti, esercenti, artigiani: coloro che gestiscono, in conto proprio, un negozio o un esercizio pubblico (bar, ecc.). Dello stesso gruppo fanno parte gli agenti di commercio, che esercitano in modo autonomo (senza dipendere da un'azienda) un'attività commerciale, ma senza avere un'azienda commerciale propria.

Artigiani sono coloro che gestiscono, in conto proprio, un'azienda artigiana che ha meno di 10 dipendenti (altrimenti vengono classificati fra gli imprenditori).

- (9) Altri lavoratori in proprio: coloro che esercitano un'attività in modo autonomo senza rapporto di dipendenza né vincolo di subordinazione e senza avere un'azienda (agenti di commercio, rappresentanti, infermieri, imbianchini, riparatori, tassisti, fotografi, ambulanti, facchini, custodi e simili).
- (10) Agricoltori: conduttori proprietari o affittuari e simili del fondo. La stessa qualifica viene data ai familiari coadiuvanti.
- (11) Insegnanti: maestri elementari, professori di scuola media inferiore e superiore, docenti universitari.
- (12) Giornalisti, artisti: liberi professionisti o dipendenti.
- (13) Operai: coloro che esercitano, contro retribuzione, una funzione esecutiva con prevalente lavoro manuale, nell'industria, nel terziario o nella pubblica amministrazione.
- (14) Agricoltori dipendenti: operai nel settore agricolo, cioè salariati dei coltivatori diretti o di aziende agricole. Sono compresi anche i salariati della pesca, silvicoltura e allevamenti.

I non occupati vengono classificati nelle seguenti "categorie non professionali":

- (15) Casalinghe: donne che vivono con reddito del marito o di altri familiari.
- (16) Studenti: persone che frequentano scuole di ogni ordine e grado purché non abbiano un'occupazione (gli studenti-lavoratori sono attribuiti ai gruppi professionali corrispondenti alla natura del lavoro svolto).
- (17) Pensionati: uomini e donne che non lavorano e che godono di una rendita fissa mensile erogata da un Istituto di previdenza, di pensione di reversibilità o di pensione sociale o di invalidità.
- (18) Altri: comprendono i disoccupati, le persone in cerca di prima occupazione e altri gruppi non altrimenti classificabili (religiosi, militari, lavoratori a domicilio, familiari coadiuvanti).

Categoria socio-professionale

Il carattere "categoria socio-professionale" deriva da vari raggruppamenti o incroci dei gruppi professionali e non professionali elencati nel paragrafo precedente.

Nel seguente prospetto è definita l'origine della classificazione secondo categorie socio-professionali.

DEFINIZIONE DEL CARATTERE "CATEGORIA SOCIO-PROFESSIONALE"

<i>Definizione</i>	<i>Quali gruppi sono compresi (professione dell'intervistato o del capofamiglia)</i>
<i>Ceti superiori</i>	<i>Imprenditori, liberi professionisti, dirigenti, quadri, ufficiali superiori e possidenti</i>

<i>Ceti medi</i>	<i>Impiegati, commercianti, artigiani, coadiuvanti, lavoratori in proprio, militari e religiosi</i>
<i>Agricoltori</i>	<i>Agricoltori conduttori e coadiuvanti</i>
<i>Intellettuali docenti</i>	<i>Insegnanti, giornalisti, artisti</i>
<i>Intellettuali studenti</i>	<i>Studenti</i>
<i>Operai</i>	<i>Operai</i>
<i>Braccianti</i>	<i>Salariati agricoli</i>
<i>Pensionati e altri</i>	<i>Pensionati e altri non occupati</i>
<i>Casalinghe di famiglie non operaie</i>	<i>Quando il capofamiglia appartiene a gruppi diversi da "operaio" e "salariato agricolo"</i>
<i>Casalinghe di famiglie operaie</i>	<i>Quando il capofamiglia appartiene ai gruppi "operaio" o "salariato agricolo"</i>

Condizione professionale

Il carattere deriva da raggruppamenti di condizioni professionali e non professionali talvolta incrociate con la condizione professionale del capofamiglia

Nel seguente prospetto è definita l'origine della classificazione.

DEFINIZIONE DEL CARATTERE "CONDIZIONE PROFESSIONALE"

Condizione (Codice riportato nel nastro di pianificazione)	Quali gruppi sono compresi (professione dell'intervistato e del capofamiglia)
1	<i>Imprenditori, dirigenti, alti funzionari, liberi professionisti, proprietari, redditieri, benestanti</i>
2	<i>Impiegati o categorie intermedie</i>
3	<i>Negozianti, esercenti, artigiani con azienda</i>
4	<i>Agenti di commercio, rappresentanti (autonomi), altri lavoratori in proprio senza azienda</i>
5	<i>Agricoltori conduttori, familiari coadiuvanti di agricoltori</i>
6	<i>Insegnanti, giornalisti, artisti</i>

7	<i>Operai (o assimilati), agricoltori dipendenti, braccianti</i>
8	<i>Casalinghe con capofamiglia di cat. 1</i>
9	<i>Casalinghe con capofamiglia di cat. 2, 6</i>
10	<i>Casalinghe con capofamiglia di cat. 3, 4, 5</i>
11	<i>Casalinghe con capofamiglia di cat. 7</i>
12	<i>Casalinghe con capofamiglia di altre categorie</i>
13	<i>Studenti</i>
14	<i>Pensionati</i>
15	<i>Militari e paramilitari, religiosi, familiari coadiuvanti, lavoratori a domicilio o loro coadiuvanti, in cerca di prima occupazione, disoccupati, altro</i>

Avvertenze per la lettura dei dati

Variabili utilizzate per le stime di lettura

Le stime di lettura "GIORNO MEDIO" per i quotidiani e "ULTIMO PERIODO" per i periodici sono basate sulle dichiarazioni degli intervistati.

La forma dei dati

I dati vengono espressi, a seconda dei casi, in una o più delle seguenti forme:

Valori assoluti: stime in '000

Valutazione dell'universo corrispondente, in migliaia di persone.

Il totale dell'universo rappresentato dal campione è pari a circa 53.000, che significa 53.000.000 di individui di 13 anni e oltre.

Percentuali di penetrazione

Percentuali calcolate sul totale dell'universo (oppure su sottogruppi significativi di esso: per es. tutti gli uomini, tutte le donne, tutti i giovani da 13 a 24 anni; ecc.), per indicare la penetrazione del fenomeno di lettura oggetto di studio. Per es. su 100 adulti italiani - oppure su 100 uomini, su 100 donne residenti in Italia - quanti sono lettori della testata X.

Percentuali di composizione

Percentuali calcolate sul totale dei lettori di una data testata, per indicare la composizione di tale collettività secondo vari caratteri di segmentazione (come il sesso, l'età, la classe sociale, il titolo di studio, ecc.). Per es. su 100 lettori della testata X quanti hanno da 13 a 17 anni.

Arrotondamenti

A causa di arrotondamenti inevitabili la somma di due o più stime o percentuali di composizione (p. es. quella dei lettori maschi e quella delle lettrici di una testata) può differire leggermente dal totale atteso (p. es. totale lettori adulti della testata).

Al fine di rendere la stima più esatta, nei pesi utilizzati per il calcolo delle tavole sono utilizzati 6 decimali.

Numero di contatti per copia (Quotidiani)

Ai lettori ultimi 7 giorni delle singole testate viene posta la domanda per valutare il numero di volte in cui la testata è stata presa in mano nell'ultimo giorno di lettura, per leggerla o sfoglarla.

Da questa informazione viene ricavato "il numero di contatti di lettura per copia".

Questa stima è statisticamente valida in quanto le interviste sono ben distribuite rispetto ai giorni della settimana, dando una buona copertura per i 7 giorni di uscita.

Pertanto, l'ultimo contatto descritto da ogni intervistato per ogni testata letta o sfogliata rappresenta bene il "contatto medio".

Numero di contatti per pagina (Periodici)

Ai lettori ultimo periodo delle singole testate vengono poste domande per valutare:

V - il numero di volte in cui la testata (non importa quale o quali numeri) era stata presa in mano, per leggerla o sfoglarla, nell'ultimo periodo;

P - la percentuale di pagine lette o sfogliate durante l'ultima occasione in cui ha letto o sfogliato la rivista (ultimo contatto).

Da queste due informazioni viene ricavato il dato "numero di contatti di lettura per pagina" (C), in base alla seguente formula:

$$C = P \times V.$$

Questa formula è statisticamente valida, in quanto le interviste sono ben distribuite lungo l'intero arco della settimana o del mese, cioè nell'intervallo di tempo che trascorre tra l'uscita di due numeri successivi della testata (ultimo periodo). Pertanto, l'ultimo "contatto", descritto da ogni intervistato per ogni rivista letta o sfogliata nell'ultimo periodo, rappresenta bene il "contatto medio".

Analisi della Cumulazione dell'audience, ovvero lettori raggiunti dopo 2, 3, 6, 9 uscite (Quotidiani)

Lo sviluppo dell'audience raggiunta dopo 2, 3, 6, ecc. numeri è ottenuto con il metodo dei pattern di lettura individuali, contando e cumulando ad ogni uscita i valori (1=Si) all'interno del pattern fino all'uscita di riferimento.

La stima di lettori raggiunti dopo 3 mesi viene sostituita con la stima del numero lettori ultimi 3 mesi basata sulle dichiarazioni degli intervistati, che coincide con il totale dei

lettori esposti ad almeno 1 uscita su 29. La stima dei lettori "giorno medio" è basata sulle dichiarazioni degli intervistati.

Analisi della Cumulazione dell'audience, ovvero lettori raggiunti dopo 2, 3, 6, 9 uscite (Periodici)

Lo sviluppo dell'audience raggiunta dopo 2, 3, 6, ecc. numeri è ottenuto con il metodo dei pattern di lettura individuali, contando e cumulando ad ogni uscita i valori (1=Si) all'interno del pattern fino a all'uscita di riferimento.

La stima dei lettori raggiunti dopo 12 uscite coincide con la stima del numero dei lettori ultimi 12 mesi (mensili) o ultimi 3 mesi (settimanali e supplementi settimanali di quotidiani), basata sulle dichiarazioni degli intervistati. La stima dei lettori "ultimo periodo" è basata sulle dichiarazioni degli intervistati.

Analisi della Probabilità di lettura per valori di frequenza dichiarata e per classi di frequenza dichiarata

I valori di probabilità esposti derivano dal rapporto tra le stime di lettura del periodo breve "B" ("ultimo periodo" per periodici e supplementi di quotidiani e "giorno medio" per i quotidiani), e quelle del periodo lungo "A" ("ultimi 12 mesi" per mensili, "ultimi 3 mesi" per i settimanali, supplementi settimanali di quotidiani e per i quotidiani), calcolato nell'ambito dei lettori delle varie classi di frequenza.

Le stime della frequenza di lettura utilizzate per l'analisi sono dedotte direttamente dalle dichiarazioni degli intervistati.

Note riguardanti determinate testate

Supplementi settimanali di quotidiani

I metodi di rilevamento adottati per i supplementi settimanali gratuiti di quotidiani e per i supplementi settimanali a pagamento di quotidiani iscritti all'indagine sono gli stessi adottati per i settimanali.

Quotidiani e supplementi di quotidiani inseriti nelle tavole regione per regione

Nelle tavole in cui i dati dei quotidiani e supplementi dei quotidiani vengono presentati regione per regione, le testate considerate in ogni singola regione sono quelle che hanno raggiunto nella regione la stima di 50 mila lettori giorno medio.

Quotidiani che escono 6 giorni su 7

Il calcolo dei "lettori giorno medio" per i quotidiani che escono 7 giorni su 7 viene fatto sommando i lettori che hanno letto nell'ultimo lunedì, nell'ultimo martedì, nell'ultimo mercoledì, ecc., fino all'ultima domenica, e dividendo la somma per 7. Il calcolo della media naturalmente deve essere modificato quando si tratta di quotidiani che escono 6 giorni su 7: per questi quotidiani, l'eventuale lettore dichiaratosi tale per il giorno di non uscita viene tolto dai lettori di quel giorno ed aggiunto ai lettori del giorno precedente, se per il giorno precedente non si era dichiarato lettore della testata, e la somma viene divisa per 6 invece che per 7.

Quotidiani che escono 5 giorni su 7

Per i quotidiani che escono 5 giorni su sette, il calcolo dei "lettori giorno medio" si basa su una divisione per cinque invece che per sette: totale lettori 1° giorno + totale lettori 2° giorno + ... + totale lettori 5° giorno (esclusi i lettori dei giorni di non uscita), diviso cinque.

AVVENIRE (Quotidiano)

Precisazione su richiesta dell'Editore di "AVVENIRE":

Nell'interpretare i dati relativi sia al numero che al profilo dei lettori di "AVVENIRE", occorre ricordare che dal campione sono esclusi i membri delle convivenze e comunità religiose, che rappresentano, come è noto, una quota rilevante degli acquirenti della testata, ed in particolare degli abbonati, come risulta dalla dichiarazione dell'Editore all'ADS.

Probabilità e frequenza di lettura

Viene pubblicata la frequenza di lettura per singolo valore di frequenza dichiarata. Oltre alla frequenza dettagliata è pubblicata anche la frequenza di lettura in classi. La definizione utilizzata per le classi di lettura riportate nelle tavole "Valore medio della probabilità individuale di lettura per classi di frequenza dichiarata" è la seguente:

Classe di frequenza	Frequenza dichiarata da questionario	
	Quotidiani	Periodici e supplementi di quotidiani
<i>bassa</i>	<i>meno di 1 giorno alla settimana</i>	<i>da 1 a 3 numeri su 12</i>
<i>media</i>	<i>da 1 giorno alla settimana a 3 giorni alla settimana</i>	<i>da 4 a 8 numeri su 12</i>
<i>alta</i>	<i>da 4 giorni alla settimana a 7 giorni alla settimana</i>	<i>alta da 9 a 12 numeri su 12</i>

La stima di lettura "ultimo periodo", riportata nelle tavole "Valore medio della probabilità individuale di lettura per valori di frequenza dichiarata e per classe di frequenza dichiarata" relativa ai periodici e supplementi di quotidiani, viene ottenuta in modo identico alle stime di lettura "ultimo periodo", utilizzando la dichiarazione di "ultima lettura" dell'intervistato.

9. AUDICOM DATABASE RESPONDENT LEVEL DIGITAL & PRINT

9.1. Fusione rilevazione digital panel e stampa

L'obiettivo della metodologia è rendere disponibile una misurazione congiunta tra la rilevazione print e quella digital. L'obiettivo finale è quello di consentire una stima congiunta delle due rilevazioni, con lo scopo di valutare la deduplicazione tra le entità misurate nella ricerca digitale e in quella stampa.

Per tali rilevazioni saranno utilizzati due panel separati. Per tale motivo, serve ricorrere a modelli statistici, in questo caso la data fusion, come approccio metodologico consolidato per unificare le misurazioni provenienti da campioni diversi. Tale metodo viene comunemente utilizzato per integrare set di dati utilizzando analisi statistiche e modelli matematici al fine di creare un singolo dataset che incorpora gli attributi di entrambi

9.2. Metodologia per il collegamento delle due ricerche

La fusione tra due panel fornisce un collegamento tra i record di due diversi campioni.

Da questo collegamento, possiamo assegnare le variabili di un record del panel “donatore” al record corrispondente del panel “ricevente”. Queste variabili sono variabili fuse.

L'obiettivo è preservare la correlazione e la coerenza tra le variabili fuse nonché con le variabili del campione ricevente.

Per effettuare questa operazione è necessario disporre di variabili comuni tra i due campioni (ad esempio variabili demografiche, variabili di abitudine dei media, ecc.). Queste variabili comuni sono utilizzate come variabili di collegamento. Viene applicato un metodo avanzato per massimizzare l'utilizzo delle variabili di collegamento, infatti i record sono collegati in base alla minima distanza delle variabili tra i panelisti. Inoltre, verranno attivate statistiche diagnostiche finali per valutare la qualità della fusione.

L'approccio usato è quello di collegare le ricerche includendo la rilevazione stampa in quella digitale. Quest'ultima, infatti, dispone di misurazione continuativa e passiva meterizzata, su campioni estesi, con frequenza di rilascio e di misurazione superiore all'indagine per la stampa.

Il collegamento statistico tra le due currency, si fonda su alcune considerazioni:

- Il collegamento sarà eseguito a livello di singolo record (record level), simulando una misurazione single source sulla totalità dei campioni;

- Il collegamento sarà effettuato sulla parte di popolazione comune; le parti non comuni saranno riportate al di fuori del processo di fusione (ad es. fruitori stampa senza accesso al digitale).

9.3. Data Fusion

Il modello di fusione è stato disegnato appositamente per collegare due (o più) campioni, a record level. Tipicamente la fusione collega i record di un campione 'donatore', ai record di un campione 'ricevente'. Come accennato, il campione "ricevente" sarà la rilevazione digitale, quello "donatore" la ricerca stampa, per le ragioni precedentemente esposte. In questo modo la ricerca digitale resterà in principio invariata, potendo continuare ad avvalersi dell'applicazione del processo sRLD precedentemente descritto.

Il collegamento utilizzerà informazioni comuni armonizzate provenienti:

- dalla componente single source che ci si attende a valle delle sinergie che verranno messi in campo con le attività di cross-recruitment, che a sua volta sarà estesa tramite specifica modellizzazione
- Da informazioni aggiuntive rilevate nel questionario stampa (comportamento di fruizione digitale)

Un elemento importante per la qualità del collegamento tra i record delle diverse rilevazioni è la disponibilità di informazioni comuni ai due campioni, in particolare quelle utili a calcolare la duplicazione. Ciò verrà ottenuto tramite i mediahook, ovvero la rilevazione delle abitudini di consumo dei media digitali in entrambi i campioni (Panel Digitale e intervista Stampa) in modo coerente e comparabile.

Il processo di fusione viene effettuato mensilmente, in modo da poter utilizzare il dato stampa (cumulato) con la stessa frequenza di quella Digitale (mensile). All'interno del periodo di validità della rilevazione stampa la ricerca donatrice resterà costante, fatto salvo le normali variazioni della popolazione digitale e le fluttuazioni del campione (turnover mensile).

Il risultato finale del processo di fusione digital + print confluirà nel Database Respondent Level Digital+Print.

Tecnicamente tale database sarà declinato in diversi file:

- A. file di descrizione delle attività di navigazione web
- B. file di descrizione delle attività di lettura stampa (ad-hoc mensile)
- C. linkage file per l'armonizzazione e la deduplica tra i file Digital e Print
- D. file demografico

10. CATALOGAZIONE DELL'OFFERTA EDITORIALE E PUBBLICITARIA

10.1. Costruzione del MarketView

Le regole di classificazione e la gerarchia di reporting hanno lo scopo di fornire una maggiore granularità utilizzando le regole e la struttura qui riportate qui di seguito.

Il sistema di classificazione MarketView viene utilizzato anche per classificare i volumi dei consumi rilevati nella soluzione "Audicom ADV Census" allineandoli al sistema di classificazione editoriale.

Livello Parent

Il Parent è il più alto livello di aggregazione contiene il roll-up di tutti i livelli sottostanti della gerarchia. Il traffico proveniente da questi livelli inferiori costituisce il totale complessivo riportato a livello Parent.

Il Parent è definito come: una organizzazione che possiede un numero di azioni con diritto di voto sufficiente per controllare la gestione e le operazioni di un'altra entità, esercitando un'influenza o eleggendo il suo consiglio di amministrazione. L'organizzazione "Parent" rappresenta il consolidamento di un gruppo di domini e di URL che sono di proprietà di una specifica società, delle sue subsidiarie o unità operative. Inoltre, un Parent può essere anche rappresentato da un'organizzazione, ente governativo, gruppo privato, società o altra istituzione, che ha partecipazioni di controllo in ogni dominio e URL del gruppo.

Il livello Parent sarà disponibile nel report mensile.

Livello Brand

Questo livello è costituito dall'aggregazione di tutte le entità al di sotto di essa, essendo il livello sottostante definito come Sub-Brand. Il traffico proveniente da questi livelli più bassi contribuirà al totale di traffico per il livello Brand. Questo traffico complessivo a livello di Brand, a sua volta, contribuirà al traffico totale per il suo Parent di riferimento:

- un Brand deve rappresentare un nome identificativo che distingue un prodotto o azienda dai suoi concorrenti.
- un Brand deve essere chiaramente riportato e utilizzato anche per il livello sottostante (Sub-Brand). Il livello Brand è disponibile in tutti i report.

Livello Sub-Brand

Il livello Sub-Brand ha lo scopo di rappresentare un raggruppamento di entità o piattaforme all'interno di uno specifico Brand così come viene offerto all'utenza finale. Come tale può essere, ad esempio, la combinazione di sito web e mobile App appartenenti allo stesso Brand.

L'obiettivo finale è quello di fornire audience deduplicate su più piattaforme in modo che i sottoscrittori possano effettuare una lettura accurata della portata effettiva delle loro Property. Il livello Sub-Brand risiede sempre sotto il livello Brand. Il roll-up di questo livello fornisce il totale del livello Brand.

Il livello Sub-Brand è disponibile in tutti i report.

Categorie e Sub-Categorie

Il MarketView prevede la classificazione delle entità anche per categorie e sub-categorie. Lo scopo delle categorie è quello di raccogliere domini e URL sulla base di conformità e analogia dei contenuti. Sono disponibili 15 Categorie che contengono un totale di 85 sottocategorie.

Categoria	Subcategoria
<i>Automotive</i>	<i>Automotive Information</i>
<i>Automotive</i>	<i>Automotive Parts, Accessories & Services</i>
<i>Automotive</i>	<i>Multi-category Automotive</i>
<i>Computers & Consumer Electronics</i>	<i>Software Information / Developers</i>
<i>Computers & Consumer Electronics</i>	<i>News & Information</i>
<i>Computers & Consumer Electronics</i>	<i>Hardware Manufacturers</i>
<i>Computers & Consumer Electronics</i>	<i>Photography Equipment & Services</i>
<i>Computers & Consumer Electronics</i>	<i>Multi-category Computers & Consumer Electronics</i>
<i>Corporate & Company Info</i>	<i>Corporate & Company Info</i>
<i>Education & Careers</i>	<i>Career Development</i>
<i>Education & Careers</i>	<i>Educational Resources</i>
<i>Education & Careers</i>	<i>Multi-Category Education & Careers</i>
<i>Education & Careers</i>	<i>Universities</i>

<i>Entertainment</i>	<i>Broadcast Media</i>
<i>Entertainment</i>	<i>Humor</i>
<i>Entertainment</i>	<i>Multi-category Entertainment</i>
<i>Entertainment</i>	<i>Adult</i>
<i>Entertainment</i>	<i>Online Games (Video/PC Gaming)</i>
<i>Entertainment</i>	<i>Videos/Movies</i>
<i>Entertainment</i>	<i>Sports</i>
<i>Entertainment</i>	<i>Gambling/Sweepstakes</i>
<i>Entertainment</i>	<i>Arts/Graphics</i>
<i>Entertainment</i>	<i>Music</i>
<i>Entertainment</i>	<i>Events</i>
Categoria	Subcategoria
<i>Entertainment</i>	<i>Books</i>
<i>Family & Lifestyles</i>	<i>Pets</i>
<i>Family & Lifestyles</i>	<i>Religion & Spirituality</i>
<i>Family & Lifestyles</i>	<i>Family Resources</i>
<i>Family & Lifestyles</i>	<i>Kids, Games, Toys</i>
<i>Family & Lifestyles</i>	<i>Genealogy</i>
<i>Family & Lifestyles</i>	<i>Personals</i>
<i>Family & Lifestyles</i>	<i>Multi-category Family & Lifestyles</i>
<i>Family & Lifestyles</i>	<i>Health, Fitness & Nutrition</i>
<i>Finance/Insurance/Investment</i>	<i>Credit Card</i>
<i>Finance/Insurance/Investment</i>	<i>Online Trading</i>
<i>Finance/Insurance/Investment</i>	<i>Financial News & Information</i>
<i>Finance/Insurance/Investment</i>	<i>Multi-category Finance/Insurance/Investments</i>

<i>Finance/Insurance/Investment</i>	<i>Full Service Banks & Credit Unions</i>
<i>Finance/Insurance/Investment</i>	<i>Financial Tools</i>
<i>Finance/Insurance/Investment</i>	<i>Loans</i>
<i>Finance/Insurance/Investment</i>	<i>Insurance</i>
<i>Government & Non-Profit</i>	<i>Military</i>
<i>Government & Non-Profit</i>	<i>Government</i>
<i>Government & Non-Profit</i>	<i>Non-Profit</i>
<i>Home & Fashion</i>	<i>Apparel/Beauty</i>
<i>Home & Fashion</i>	<i>Multi-category Home & Fashion</i>
<i>Home & Fashion</i>	<i>Cooking, Food & Beverages</i>
<i>Home & Fashion</i>	<i>Home & Garden</i>
<i>Home & Fashion</i>	<i>Real Estate/Apartments</i>
<i>Multi-category Commerce</i>	<i>Shopping Directories & Guides</i>
<i>Multi-category Commerce</i>	<i>Classifieds/Auctions</i>
<i>Multi-category Commerce</i>	<i>Free Merchandise</i>
<i>Multi-category Commerce</i>	<i>Coupons/Rewards</i>
<i>Multi-category Commerce</i>	<i>Multi-category Retail</i>
<i>News & Information</i>	<i>Current Events & Global News</i>
<i>News & Information</i>	<i>Directories/Local Guides</i>
<i>News & Information</i>	<i>Multi-category News & Information</i>
<i>News & Information</i>	<i>Special Interest News</i>
<i>News & Information</i>	<i>Weather</i>
<i>News & Information</i>	<i>Research Tools</i>
<i>Search Engines/ Communities/ Social Networking</i>	<i>Search</i>

<i>Search Engines/ Communities/ Social Networking</i>	<i>Targeted Portals & Communities</i>
<i>Search Engines/ Communities/ Social Networking</i>	<i>General Interest Portals & Communities</i>
<i>Search Engines/ Communities/ Social Networking</i>	<i>Social Networks & Member Communities</i>
<i>Special Occasions</i>	<i>Greeting Cards</i>
<i>Special Occasions</i>	<i>Gifts & Flowers</i>
<i>Special Occasions</i>	<i>Delivery/Stamps</i>
<i>Special Occasions</i>	<i>Holidays & Special Events</i>
<i>Special Occasions</i>	<i>Multi-category Special Occasions</i>
<i>Telecom/Internet Services</i>	<i>Cellular/Mobile</i>
<i>Telecom/Internet Services</i>	<i>Multi-category Telecom/Internet Services</i>
<i>Telecom/Internet Services</i>	<i>Web Hosting Services</i>
<i>Telecom/Internet Services</i>	<i>Internet Tools/Web Services</i>
<i>Telecom/Internet Services</i>	<i>ISP</i>
<i>Telecom/Internet Services</i>	<i>Long Distance/Local Carrier</i>
<i>Telecom/Internet Services</i>	<i>Instant Messaging</i>
<i>Telecom/Internet Services</i>	<i>E-mail</i>
<i>Travel</i>	<i>Ground Transportation</i>
<i>Travel</i>	<i>Travel Destination & Attractions</i>
<i>Travel</i>	<i>Maps/Travel Info</i>
<i>Travel</i>	<i>Multi-category Travel</i>
<i>Travel</i>	<i>Hotels/Hotel Directories</i>
<i>Travel</i>	<i>Airlines</i>
<i>Travel</i>	<i>Cruise Lines</i>

Attribuzione a Sub-Brand tematici dei consumi di contenuti editoriali statici / testuali fruiti mediante Google AMP

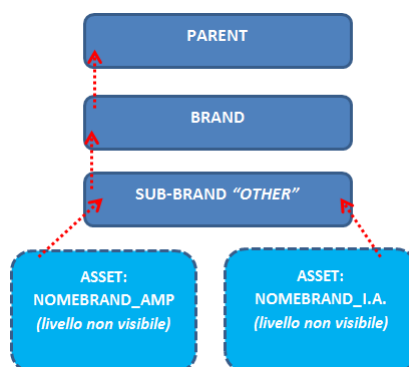
Il servizio in oggetto è a pagamento ed è attivabile mediante sottoscrizione da tutti i Publisher iscritti. È altresì già prevista ed applicata una modalità di catalogazione dei contenuti fruiti mediante Google AMP che non prevede costi aggiuntivi, descritta qui di seguito come “modalità standard”.

Posto che tali modalità di distribuzione dei contenuti siano già implementate dal Publisher, l'unico requisito richiesto per attivare la modalità di attribuzione “standard” è avere completato i processi di certificazione dell'installazione di SDK Nielsen per Google AMP.

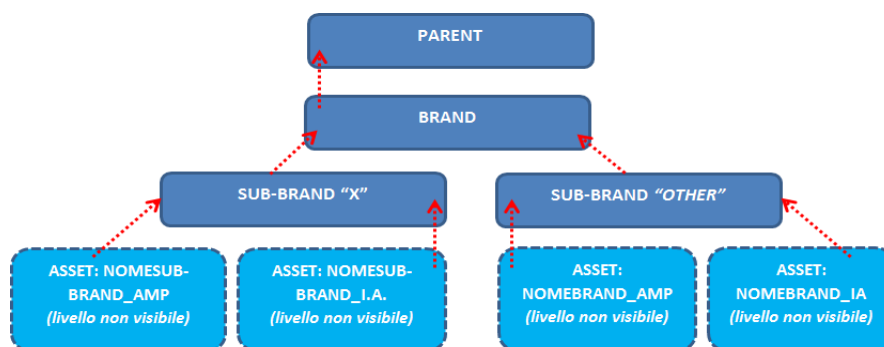
I Publisher iscritti che vogliono aderire al servizio che consente la modalità di attribuzione nei Sub-Brand devono procedere per ogni piattaforma di distribuzione (AMP) con nuove certificazioni per ogni Sub-Brand (diverso da quello “Other/Altro”) in cui far confluire anche i consumi delle piattaforme stesse.

Modalità di attribuzione “standard”:

Il totale dei consumi dei contenuti di un Brand distribuiti via Google AMP viene attribuito attraverso la compilazione di appositi parametri di SDK da parte del Publisher in due “Asset” ad-hoc che riportano gerarchicamente al Sub-Brand “Other/Altro” e quindi al totale del “Brand”.



Modalità di attribuzione nei Sub-Brand:



In questa modalità anche i consumi del “Sub-Brand X” mediante Google AMP sono attribuiti al “Sub-Brand X” e quindi riportati al “Brand”. Altri eventuali consumi del Brand via Google AMP non incanalati nel Sub-Brand tematico continueranno ad essere attribuiti al “Sub-Brand – Other/Altro” e quindi al totale del Brand.

I Publisher iscritti che aderiscono al servizio devono procedere per ogni piattaforma di distribuzione (AMP) con nuove certificazioni per ogni Sub-Brand (diverso da quello “Other/Altro”) in cui far confluire anche i consumi delle piattaforme stesse.

I livelli “Asset” concorrono alla produzione del dato totale del Sub-Brand ma nel Sub-Brand non sarà possibile visualizzarli nel dettaglio.

L’attivazione del servizio diventa tassativa per i Publisher iscritti che decidono di procedere con la rilevazione dei consumi mediante Google AMP mediante certificazione di “Sub-Brand TAL” al fine di poter attribuire coerentemente i consumi organici rispetto a quelli non-organici.

Per i Publisher iscritti dotati di Sub-Brand TAL che siano provvisti di SDK che consente la rilevazione del traffico AMP è necessario l’adeguamento alla regola di attribuzione.

10.2. Traffic Assignment Letter (TAL)

Per la rilevazione dei dati delle audience digitali distribuiti da Audicom adotta un sistema di regole per la gestione del “Catalogo” dell’offerta editoriale su Internet, strutturato per gerarchia di navigazione (Parent, Brand, Channel / Sub-Brand...) e organizzato per categorie di contenuti editoriali e per macro-aggregazioni. Il sistema di classificazione riflette la necessità di rappresentare in modo totalmente trasparente le componenti di audience “organiche” da quelle “aggregate” dove per audience “organiche” si intendono quelle derivanti dal traffico realizzato su siti di proprietà del publisher iscritto (Cessionario) mentre per “aggregate” si intendono quelle derivanti da accordi di cessione del traffico da parte di altro publisher (Cedente).

Gli accordi sulla cessione del traffico erano originariamente governati da Nielsen per conto di Audicom Srl mediante la sottoscrizione di un documento denominato “T.A.L.” (Traffic Assignment Letter).

Elenco delle regole

1. la T.A.L. è un documento gestito da Audicom e non più in autonomia da Nielsen Media Italy per tutti gli editori che pubblicano nel nastro (Audicom Database Respondent Level Digital). Rimane invece nella sua forma attuale per altri siti presenti solo nel report mensile aggregato (disponibile in Audicom Media View).
2. le T.A.L. in ambito Audicom Database Respondent Level Digital vengono governate direttamente da Audicom che si avvale di Nielsen Media Italy per il solo servizio tecnico di gestione. Per dirimere i casi incerti provvede autonomamente il team di

gestione della società: in caso di difficoltà o per ricorso di terzi, la questione viene sottoposta al Comitato Tecnico con successiva ratifica consiliare;

3. il logo del Brand cessionario deve essere riportato chiaramente e con evidenza nelle pagine dei siti Sub-Brand, in modo che questi risultino come Sezioni o Componenti del Brand. È preferibile, anche se non obbligatorio, l'utilizzo dei sub-domain.
4. il logo del Brand cessionario deve essere riportato in tutte le pagine del Brand cedente all'interno di uno slimheader, insieme ad un testo di accompagnamento. Due i template a disposizione: il primo con slimheader mobile, cioè che scorre insieme alla pagina durante lo scrolling (di seguito "Template 1"); il secondo con slimheader fisso, cioè che resta sempre visibile anche durante lo scrolling della pagina (di seguito "Template 2"). I due template presentano alcune caratteristiche comuni ed alcune esclusive

Di seguito sono riportate tutte le indicazioni per una adeguata conformazione del look & feel del Brand cedente. Tali indicazioni per il momento trovano applicazione per il solo mondo desktop, mentre sono allo studio applicabilità ed eventuale adattabilità all'ambito Mobile (browsing e App).

CARATTERISTICA	TEMPLATE 1	TEMPLATE 2
Tipologia <i>slimheader</i>	Mobile scorre insieme alla pagina	Fisso sempre visibile in posizione fissa
Posizionamento <i>slimheader</i>	Sopra <i>header</i> Sito	Sempre fisso all'inizio della pagina
Altezza minima <i>slimheader</i>	45 pixel	40 pixel
Larghezza <i>slimheader</i>	Come <i>header</i> del sito Cedente sottostante	Al 100%
Colore sfondo <i>slimheader</i>	Libero	
Allineamento del logo e del testo di accompagnamento	A sinistra o al centro	
Logo	Medesimo logo riportato nella homepage del Brand Cessionario (dal quale può essere tolto il <i>pay off</i>)	
Logo Cliccabile	SI	

Altezza minima Logo	40 pixel	35 pixel
Distanza del logo sito Cessionario dal logo del sito Cedente	Non possono esserci elementi pubblicitari, ma solo contenuti editoriali	Non ci sono vincoli
Testo di accompagnamento	"Questo sito contribuisce alla audience di [logo Brand]"	
Font	Family e color liberi	
Grandezza minima Font	14 pixel	13 pixel

5. qualora il Brand aggreghi Sub-Brand che si configurino come componenti, ovvero siti autonomi frutto di cessioni di traffico (T.A.L.), il Brand deve avere un nome/logo chiaramente differente da quello di ogni Sub-Brand che la compone ed ogni SubBrand deve avere un nome o logo riportato nelle pagine ricomprese;
6. tutte le cessioni oggetto di T.A.L. devono comparire in Sub-Brand diversi e separati da quelli organici ed essere evidenziate come non organiche in tutta la reportistica, che dovrà evidenziare, in una apposita nota, anche gli editori conferenti, ossia:
 - a. i siti in T.A.L. saranno separati dai siti «organici», che mantengono il nome Sub-Brand come originariamente codificato;
 - b. il *naming* per tutti i raggruppamenti di T.A.L. è:
 - i. NOME BRAND + "TAL" + NOME CUSTOM(*) + SUB CAT (in cui è codificato il Sub-Brand attuale che lo contiene);
 - ii. le TAL appartenenti a sottocategorie omogenee sono raggruppate dentro uno stesso Sub-Brand.

(*) La compilazione del campo «Nome Custom» da parte del publisher titolare del Sub-Brand è facoltativa.

10.3. Collegamento testate "Print" con rilevazione "Digital" per distribuzione dati "Digital&Print"

L'impianto della Ricerca Integrata necessita della convergenza dei cataloghi utilizzati per la rilevazione digitale e stampa. Ciò è fondamentale non solo per garantire che i dati distribuiti possano essere ricondotti agli stessi oggetti editoriali, ma anche per rendere possibile l'applicazione del modello di fusione.

Al tal scopo, ogni testata cartacea iscritta alla rilevazione "print" deve disporre di uno specifico sub-brand con la seguente tassonomia "NOME SUB-BRAND_print". Il sub-brand verrà creato all'interno della brand digital della testata. Eccezioni: in caso di "Brand" digital cui sono collegabili più testate cartacee, per ognuna di esse è necessario creare uno specifico sub-brand con la seguente tassonomia "NOME SUB-BRAND_print".

11. ALLEGATI

11.1. Schema del Database Respondent Level Digital

Ogni nome del file contiene il periodo di validità di produzione, nel formato YYMM.

I dati sono distribuiti mensilmente e si riferiscono esclusivamente ai dati relativi al periodo indicato con formato YYMM riportato.

Ogni campo è separato da punto e virgola “;”. Gli header dei file sono riportati nella prima riga.

Lista dei file:

1. Demographic DB: AWDB_MMY_Y_DISTR_DEMOG_SRLDa.dat,
2. AWDB_MMY_Y_SampleDesc_SRLDa.dat
3. Navigational Dataset³: AWDB_MMY_Y_STDFMT_NAVI_SRLDa.DAT
4. Dictionaries: see related chapter

I file Demographic DB e and Navigational Dataset sono ordinati per pID..

Lista completa dei file forniti mensilmente:

<i>File name</i>	<i>Content</i>	<i>Type of file</i>
AWDB_MMY_Y_DISTR_DEMOG_SRLDa.dat	Demographics data	not static
AWDB_MMY_Y_SampleDesc_SRLDa.dat	Demographics codes description	static
AWDB_MMY_Y_STDFMT_NAVI_SRLDa.dat	Navigation (viewing data)	not static
DIZ_PARENT_ACTIVE_MMY_Y_DISTRa.dat	Parent	not static
DIZ_BRAND_ACTIVE_MMY_Y_DISTRa.dat	Brand	not static
DIZ_CHANNEL_ACTIVE_MMY_Y_DISTRa.dat	Channel	not static
DIZ_APPLICATION_ACTIVE_MMY_Y_DISTRa.dat	Application	static
DIZ_OBJECT_ACTIVE_MMY_Y_DISTRa.dat	Object	static

³ Il navigational Dataset conterrà il file specifico per la fruizione Print, si veda paragrafo seguente, per comporre il Database TL combinato Digital+Print.

DIZ_DCREDIT_MMY_Y_DISTRa.dat	distributed contents	static
Categories	Categories	not static
SUBCATEGORYNAMES_MMY_Ya.csv	Subcategories	not static
CRU_MMY_Ya.csv	CRU aggregation	not static
DIZ_PIATTAFORME_ACTIVE_MMY_Y_DISTRa.dat	Device type	static
AWDB_MMY_Y_VALIDATEa.csv	Validation	not static
NOAPPa.csv	Validation	not static

Laddove

- *MMYY* = month, year: Esempio: Febbraio 2017=0217
- Static= il file non cambia mese dopo mese, eventuali eccezioni saranno specificate. Not Static= i file cambiano mese dopo mese.

Tutti i Brand e Sub-Brand che non sono sottoscrittori di di Audicom sono raggruppati nel Parent/Brand “Altro” con codice “9999998”.

Demographic DB

- Il file ASCII contiene le informazioni demografiche per ogni record.
- Le tipologie di Field sono: Numerico “N”, Testo “T” o Data “D”.
- La somma dei pesi corrisponde all’universo rappresentato (popolazione totale, inclusiva di utenti internet e non utenti internet).

AWDB_MMY_Y_DISTR_DEMOG_SRLDa.dat:

Field	ID of Demo Var	Type	Name	Description	Note
1		N	pID	record ID	
2		N	WEIGHT	Weight of panelist	Format: xxxxx.xxx, always 3 decimal digits
3		N	EMPTY	Reserved for future use	---
4	1	N	ETA	Age in years	Quantitative
5	2	N	SESSO	Gender	Categorical
6	3	N	EDUCAZIONE	Education level	Categorical

7	4	N	CONDLAV	Job situation	Categorical
8	5	N	OCCUP	Occupation	Categorical
9	6	N	COMP_2_11	count of people in hh age 2-11	Quantitative
10	7	N	COMP_12_17	count of people in hh age 12-17	Quantitative
11	8	N	COMP_TOT	household size	Quantitative
12	9	N	REDDITO	Income	Categorical
13	10	N	ATTIVO_PC	Panelist active on PC in browsing Internet during the month	Categorical
14	11	N	AREA	Geographic Area	Categorical
15	12	N	REGIONE	Region	Categorical
16	13	N	NAZIONALITA	Nationality	Categorical
17	14	N	RA	Responsible for purchases	Categorical
18	15	N	SECOND_RA	Secondary responsible for purchases	Categorical
19	16	N	CAPOFAM	Head of Household	Categorical
20	17	N	COMP_0_1	Presence of kids in hh age 0-1	Categorical
21	18	N	ATTIVO_MOBILE	Panelist active on Mobile (smart. Or tablet) in browsing Internet during the month	Categorical

pID = progressive record ID (1 to n), not linkable per month.

AWDB_MMYT_SampleDesc_SRLD.dat:

Il file ASCII contiene le descrizioni di ogni demografiche.

Field	Type	Name	Description	Note
1	T	DOMANDA	Question text	Short version of the questionnaire text
2	N	ID VARIABLE	Variable ID	Linked with the AWDB_MMYT_DISTR_DEMOG_SRLDa file

3	N	ID_CAT	Category ID	"blank" if the variable is quantitative
4	T	LABEL_CAT	Category Label	"blank" if the variable is quantitative

Navigation Dataset

Il file ASCII contiene i record di navigazione per tutti i panelisti.

Ogni campo è separato da punto e virgola “;”.

AWDB_MMYT_STDFMT_NAVI_SRLDa.DAT:

Field	Type	Name	Description	Note
1	N	pID	record ID	Linked with the AWDB_MMYT_DISTR_DEMOG_SRLDa file
2	D	DATA	Date (DDMMAA)	calendar date
3	N	PARENT	ID Parent	Linked with diz_parent_active_MMYT_distr Dictionary
4	N	BRAND	ID Brand	Linked with diz_Brand_active_MMYT_distr Dictionary
5	N	CHANNEL	ID Channel	Linked with diz_channel_active_MMYT_distr Dictionary
6	N	IOBJ	ID Object	Linked with diz_object_active_MMYT_distr Dictionary
7	N	IAPP	ID Internet Application	Linked with diz_application_active_MMYT_distr Dictionary
8	N	PV1	Page View TimeBand 1	TimeBand 0-3
9	N	TIME1	Time Spent TimeBand 1	
10	N	PV2	Page View TimeBand 2	TimeBand 3-6
11	N	TIME2	Time Spent TimeBand 2	
12	N	PV3	Page View TimeBand 3	TimeBand 6-9
13	N	TIME3	Time Spent TimeBand 3	

14	N	PV4	Page View TimeBand 4	TimeBand 9-12
15	N	TIME4	Time Spent TimeBand 4	
16	N	PV5	Page View TimeBand 5	TimeBand 12-15
17	N	TIME5	Time Spent TimeBand 5	
18	N	PV6	Page View TimeBand 6	TimeBand 15-18
19	N	TIME6	Time Spent TimeBand 6	
20	N	PV7	Page View TimeBand 7	TimeBand 18-21
21	N	TIME7	Time Spent TimeBand 7	
22	N	PV8	Page View TimeBand 8	TimeBand 21-24
23	N	TIME8	Time Spent TimeBand 8	
24	N	PIATTAFORMA	Device	Linked with diz_piattaforme_active_MMYD_distr
25	N	DCREDIT	distributed content	Linked with diz_dccredit_MMYD_distr Dictionary
26	N	EMPTY	Reserved	

ID Channel: it is always populated.

APP:

Il dettaglio per le Applicazioni (App o Browsing) non è disponibile. Il codice ID per il campo Application (IAPP) è sempre 0 (combined).

OBJECT:

Nota: se il record è attribuito a un Video (Object, IOBJ uguale a 1):

- Il codice ID per l'oggetto (IOBJ) è sempre 1 per Object, 0 altrimenti. L'Object specifico può essere identificato rispetto al suo rispettivo Padre/Brand/Canale.
- Le metriche di Visualizzazione pagina contano il numero di volte in cui il Video è stato avviato (ad esempio Visualizzazioni video)
- Il tempo trascorso in relazione al Testo (IOBJ=0) include il tempo trascorso sull'Oggetto (IOBJ=1). Quindi per ottenere la riga totale il tempo trascorso sugli oggetti deve essere filtrato.

Inoltre, un determinato oggetto appartenente a un Brand/Sub-Brand potrebbe essere visualizzato tramite un altro Brand/Sotto-Brand.

In questo caso PV e tempo speso sono attribuiti solo al proprietario del video (se un video X viene trasmesso su player video di proprietà di A, sul sito web di A, e anche su un editore B sul sito web di B, ma sempre con il lettore video di A, tutte le metriche video vengono attribuite ad A).

DISTRIBUTED CONTENT:

Il “Distributed content” non è disponibile. Il campo ‘DCREDIT’ è sempre uguale a 0.

PIATTAFORMA:

Il campo ‘PIATTAFORMA’ indica la piattaforma di fruizione Il codice ID per piattaforma (IPIATT) è 1 per PC, 2 per Mobile e 3 per CTV

Dictionaries

In tutti i file:

- Ogni campo è separato da punto e virgola “;”
- Ogni campo ha un numero fisso di caratteri

Parent: diz_parent_active_MMYD_distrA.dat

Field	Type	Name	Description	Note
1	N	EMPTY	Progressive, reserved	
2	N	IPAR	Parent ID	Linked with the AWDB_MMYD_STDFMT_NAVI_ SRLD a file (3)
3	T	LAB_PAR	Parent Label	

Brand: diz_Brand_active_MMYD_distrA.dat

Field	Type	Name	Description	Note
1	N	EMPTY	Progressive, reserved	
2	N	BRAND	Brand ID	Linked with the AWDB_MMYD_STDFMT_NAVI_ SRLD a file (4)
3	N	CAT_ID	Category ID	Linked with CategoryNames_MMYD
4	N	SUBCAT_ID	Sub-category ID	Linked with SubcategoryNames_MMYD
5	T	LAB_BRAND	Brand Label	

Channel: diz_channel_active_MMYT_distrA.dat

Field	Type	Name	Description	Note
1	N	EMPTY	Progressive, reserved	
2	N	CHANNEL	Channel ID	Linked with the AWDB_MMYT_STDFMT_NAVI_SRLDa file (5)
3	N	CAT_ID	Category ID	Linked with CategoryNames_MMYT
4	N	SUBCAT_ID	Sub-category ID	Linked with SubcategoryNames_MMYT
5	T	LAB_CHAN	Channel Label	

Object: diz_object_active_MMYT_distrA.dat

Field	Type	Name	Description	Note
1	N	EMPTY	Progressive, reserved	
2	N	IOBJ	Object ID	Linked with the AWDB_MMYT_STDFMT_NAVI_SRLDa file (6)
3	T	LAB_OBJ	Object Label	

IOBJ : 0 = text, 1= video/obj

Application: diz_application_active_MMYT_distrA.dat

Field	Type	Name	Description	Note
1	N	EMPTY	Progressive, reserved	
2	N	IAPP	Application ID	Linked with the AWDB_MMYT_STDFMT_NAVI_SRLDa file (7)
3	T	LAB_APP	Application Label	

IAPP: 0= combined

Application: DIZ_DCREDIT_MMYT_DISTRa.DAT

Field	Type	Name	Description	Note
1	N	EMPTY	Progressive, reserved	

2	N	IDCREDIT	distributed content ID	Linked with the AWDB_MMYT_STDFMT_NAVI_SRL Da file (25)
3	T	LAB_DCREDIT	distributed content Label	

IDCREDIT: 0= NOT available

Device: diz_piattaforme_active_MMYT_distrA.dat

Field	Type	Name	Description	Note
1	N	EMPTY	Progressive, reserved	
2	N	IPIATT	Device ID	Linked with the AWDB_MMYT_STDFMT_NAVI_SRL Da file (24)
3	T	LAB_PIATT	Device Label	

IPIATT: 1=PC, 2=MOBILE

Category: CategoryNames_MMYT.csv

Field	Type	Name	Description	Note
1	N	CAT_ID	Category ID	Linked with the BRAND and CHANNEL Dictionaries file (3)
2	N	CAT_NAME	Category Label	

Category: SubcategoryNames_MMYT.csv

Field	Type	Name	Description	Note
1	N	SUBCAT_ID	Sub-category ID	Linked with the BRAND and CHANNEL Dictionaries file (4)
2	N	SUBCAT_NAME	Sub-category Label	

Groups: CRU_MMYTa.DAT

Field	Type	Name	Description	Note
-------	------	------	-------------	------

1	T	CRU_LABEL	Grouping name	
2	N	CRU_ID	Grouping ID	
3	N	BRAND	Brand ID	'blank' if the channel alone belongs to the group
4	T	CHANNEL	Channel ID	'blank' if the whole Brand belongs to the group
5	T	LAB_BRAND	Brand or Channel Label	Based on the two previous fields, the Brand or channel label
6	T	LAB_CHAN	Empty	
7	N	CRU_CAT	Category ID	CRU is attributed to a category if and only if all elements of the CRU belong to the same category
8	N	CRU_SUBCAT	Sub-Category ID	CRU is attributed to a sub-category if and only if all elements of the CRU belong to the same sub-category
9	N	PC	1: device PC=YES 0: device PC=NO	It will belong to the CRU only the navigation Brand/Channel (3 and 4 field) that do also satisfy the characteristic specified in these 3 fields. Rows 10 and 11 must contain the same value (e.g.: W: 1 and A:1 or W:0 and A:0)
10		W	1: mobile and Web=YES 0: mobile and Web=NO	
11		A	1: mobile and App=YES 0: mobile and App=NO	

Note: the file contains header as first row.

- Esempi di campi 9, 10 and 11:
- "1,0,0": CRU will only contain navigation on the Brand/Channel from PC Device (field 24 in AWDB_MMYT_STDFMT_NAVI_SRLDa.dat set to 1)
- "0,1,0": CRU will only contain navigation on the Brand/Channel from Mobile Devices (field 24 in AWDB_MMYT_STDFMT_NAVI_SRLDa.dat set to 2, 3) and Web (field 7 in AWDB_MMYT_STDFMT_NAVI_SRLDa.dat set to 0)
- "0,0,1": CRU will only contain navigation on the Brand/Channel from Mobile Devices (field 24 in AWDB_MMYT_STDFMT_NAVI_SRLDa.dat set to 2, 3) and App (field 7 in AWDB_MMYT_STDFMT_NAVI_SRLDa.dat set to 1)
- "0,1,1": CRU will only contain navigation on the Brand/Channel from Mobile Devices (field 24 in AWDB_MMYT_STDFMT_NAVI_SRLDa.dat set to 2, 3)
- "1,1,1": CRU will contain all navigations on the Brand/Channel

11.2. Schema Navigation DataSet Print

Campo	Descrizione	Note
Respondent ID "Fuso"	ID univoco relativo al rispondente	Deriva dal processo di fusione con la ricerca digitale; può essere ricollegato all'ID rispondente digitale tramite opportuni file di "linkage"
Target	Carta e/o Replica	Indicatore binario che informa se il rispondente ha visto o meno anche la versione digitale replica
Brand id (testate)	ID univoco, armonizzato con il corrispondente id digitale, della testata	
Brand name (testate)	Nome della testata	
Tipo di testata	Indicatore che informa sulla periodicità della testata (mensile, settimanale, quotidiano)	
Frequenza di lettura	Lettura di una certa testata in un periodo specifico	Lettura di numeri/giorni in un periodo. Per supplementi e settimanali: da 1 num a 12 in 3 mesi. Per mensili: da 1 num a 12 in 12 mesi. Per quotidiani: da 1 a 7 gg alla settimana, circa 1 volta al mese, 2-3 volte al mese, occasionalmente in 3 mesi
Frequenza di lettura in classi	Bassa, media e alta	
Giorni di lettura	Giorni della settimana	Presente solo se testata è quotidiano.
Lettori ultimo periodo	Indicatore binario	Solo per periodici (supplementi, settimanali e mensili)
Quante volte ha letto / sfogliato	Intero da 0 a 5+ volte	Campo popolato solo se Target indica lettore Carta. Presente solo per supplementi, settimanali e mensili
Contatti per pagina/copia	Numero contatti. Decimale.	
Lettura ultimo periodo	indica la distanza temporale dall'ultima volta che si è letto. Sei classi per i mensili e cinque per i supplementi e settimanali.	Campo popolato solo se Target indica lettore Carta e solo per i periodici (supplementi, settimanali e mensili)
Fonte di provenienza della copia	Indica la provenienza della copia	
Pattern di lettura	Campo di N posizioni impostate a 0 (non lettori) o a 1 (lettori). 29	N indica il numero di campi, specifico per ciascun tipo di testata (quotidiano, settimanale, mensile)

	posizioni per quotidiani, 12 per supplementi, settimanali e mensili	
Peso giorno medio quotidiano	Valore da utilizzare per il calcolo dei Lettori Giorno Medio di un quotidiano (si ottengono per somma del peso individuale per il peso giorno medio)	Solo per quotidiani

11.3. Audicom Print - Gli intervalli fiduciari delle stime

PREMESSA

Per la valutazione dell'affidabilità statistica delle stime esposte nei materiali di ciascuna edizione dell'indagine "Audicom Print", si propone l'adozione della formula corrispondente al limite fiduciario, nell'ipotesi di campione casuale semplice, al margine di confidenza del 95% (di seguito sono indicate le modalità di applicazione della formula, che valgono sia per il target "Lettori carta e/o replica" che per il target "Lettori carta").

INTERVALLI FIDUCIARI DELLE STIME DEI QUOTIDIANI ISCRITTI ALL'EDIZIONE PRINT

a) Intervalli fiduciari a livello nazionale

Per ciascuna edizione dell'indagine vengono forniti i limiti fiduciari delle principali stime di lettura dei quotidiani, per il target Carta e/o Replica e per il target Carta. I limiti fiduciari sono calcolati per consentire di valutare qual è il margine di confidenza entro il quale può essere accolta la stima totale, in migliaia, dei "lettori ultimi sette giorni" e quella dei "lettori nel giorno medio" di ciascun quotidiano considerato.

Il valore dell'intervallo fiduciario "lettori ultimi sette giorni" è calcolato in base alla seguente formula:

$$i.f. = \pm 1,96 \sqrt{\frac{p (100 - p)}{n}}$$

dove:

i.f. = intervallo fiduciario, al livello di confidenza del 95%

p = percentuale di penetrazione

n = numero di interviste nel campione (base proporzionale nazionale esclusi i sovracampionamenti provinciali, regionali, opzionali per il campione CAPI, e interviste considerate compliant per il campione CAWI), con riferimento alla somma delle tre rilevazioni che compongono l'edizione.

L'intervallo fiduciario relativo ai "lettori giorno medio" è ottenuto moltiplicando l'intervallo fiduciario relativo ai "lettori ultimi 7 giorni" per il rapporto

lettori giorno medio

lettori ultimi 7 giorni

ATTENZIONE!

- Per le testate a diffusione nazionale a penetrazione bassa (meno di 1-2%) e con lettorato avente spiccati connotati di élite culturale o economica, la stima dell'intervallo fiduciario va accolta con riserva, data la grande difficoltà di tenere sotto controllo, anche in un campione probabilistico, la varianza.
- Per le testate a diffusione molto concentrata in una regione o provincia, è preferibile, per valutare l'intervallo fiduciario, fare riferimento agli intervalli fiduciari delle stime regionali dei quotidiani.

Si riporta un esempio di tavola fornita per ciascuna edizione dell'indagine "Audicom Print", che contiene gli intervalli fiduciari delle stime di lettura nazionali dei quotidiani iscritti all'indagine, per il target Carta e/o replica.

Esempio Tavola

INTERVALLI FIDUCIARI DELLE STIME DI LETTURA NAZIONALE DEI QUOTIDIANI ISCRITTI ALL'INDAGINE "AUDICOM PRINT"

Lettori Carta e/o Replica

QUOTIDIANI A PAGAMENTO	Lettori U7gg		Lettori GM	
	Stim in '000	Int. Fid.	Stim in '000	Int. Fid.
NOME TESTATA	...	...	...	...
...	...	...	...	...
Universo Adulti (Nome edizione)	...			
Campione (Nome edizione)	...			

	Lettori U7gg	Lettori GM
--	--------------	------------

SUPPLEMENTI SETTIMANALI GRATUITI QUOTIDIANI	<i>Stim in '000</i>	<i>Int. Fid.</i>	<i>Stim in '000</i>	<i>Int. Fid.</i>
DI				
NOME TESTATA	...	...	...	...
...	...	...	...	...
Universo Adulti (Nome edizione)	...			
Campione (Nome edizione)	...			
SUPPLEMENTI SETTIMANALI PAGAMENTO QUOTIDIANI	<i>Lettori U7gg</i>		<i>Lettori GM</i>	
A DI	<i>Stim in '000</i>	<i>Int. Fid.</i>	<i>Stim in '000</i>	<i>Int. Fid.</i>
NOME TESTATA	...	...	...	...
...	...	...	...	...
Universo Adulti (Nome edizione)	...			
Campione (Nome edizione)	...			

b) Intervalli fiduciari a livello regionale

Per i quotidiani i limiti fiduciari sono elaborati per consentire di valutare anche i margini entro i quali possono essere accolte le stime dei lettori ultimi 7 giorni e dei lettori nel giorno medio, considerate regione per regione.

Per ogni regione e per ogni quotidiano che, avendo superato nella regione la stima di 50 mila lettori giorno medio, viene considerato nelle elaborazioni della regione stessa, nelle tavole sono riportati quattro valori: la stima in migliaia dei lettori ultimi 7 giorni; il limite fiduciario della stima, espresso in migliaia; la stima in migliaia dei lettori nel giorno medio e il corrispondente limite fiduciario, sempre in migliaia.

Il numero n di interviste eseguite nella regione sulla cui base è calcolato il livello fiduciario comprende anche le interviste aggiuntive fatte a titolo di "sovracampionamento".

Il significato dei valori calcolati e i metodi per calcolarli sono quelli stessi illustrati per gli intervalli fiduciari a livello nazionale, naturalmente con riferimento al numero di interviste effettuate nella sola regione di riferimento.

Esempio: se il quotidiano X ha una stima di 500 (mila) lettori nella regione A, e per questa stima l'intervallo fiduciario al 95% è pari a 50 (mila), ciò significa che vi sono 95 probabilità su 100 che l'errore accidentale (dovuto al fattore casuale nel campionamento) da cui può essere affetta la stima non superi l'entità di 50 (mila), in più o in meno.

In altre parole, vi è la "certezza statistica" che i lettori nel giorno medio del quotidiano X, nella regione A, non siano meno di 450 mila e non siano più di 550 mila.

Si riporta un esempio di tavole regionali che vengono fornite per ciascuna edizione dell'indagine "Audicom Print", che contiene gli intervalli fiduciari delle stime di lettura regionali dei quotidiani iscritti all'indagine, per il target Carta e/o replica.

Esempio Tavola

INTERVALLI FIDUCIARI DELLE STIME REGIONALI DEI QUOTIDIANI ISCRITTI ALL'INDAGINE "AUDICOM PRINT"

Lettori Carta e/o Replica

Nome Regione	Lettori 7gg	IF 7gg	Lettori GM	IF GM
NOME TESTATA	...	...	...	...
...	...	...	...	...
Universo Adulti (Nome edizione)	...			
Campione (Nome edizione)	...			

INTERVALLI FIDUCIARI DELLE STIME DEI PERIODICI ISCRITTI ALL'INDAGINE "AUDICOM PRINT"

Per il calcolo dell'intervallo fiduciario delle stime dei lettori di periodici si propone la seguente formula, corrispondente al margine di confidenza del 95%.

$$i.f. = \pm 1,96 \sqrt{\frac{p(100 - p)}{n}}$$

essendo:

n = numero di interviste nel campione (base proporzionale nazionale esclusi i sovracampionamenti provinciali, regionali, opzionali per il campione CAPI, e interviste considerate compliant per il campione CAWI), con riferimento alla somma delle tre rilevazioni che compongono l'edizione.

p = percentuale di penetrazione sul totale campione.

Nelle tavole seguenti riportiamo, testata per testata, l'entità dell'intervallo fiduciario, cioè dell' "errore di campionamento": quell'errore in più o in meno che, con 95 probabilità su 100, non è stato superato nel fare la stima campionaria dei "lettori dell'ultimo periodo".

AVVERTENZA

Si raccomanda che, ai fini delle valutazioni del margine di errore statistico ai livelli più disaggregati, e cioè quando le percentuali sono calcolate su un numero esiguo di casi, venga adottato il noto test empirico di Cochran, secondo il quale, quando per una percentuale si verifica che il prodotto $p \times q \times n$ è inferiore a 9 (essendo p la percentuale divisa per 100 - p. es. 0,80 se 80% - q il complemento di p a 1,00 - p. es. 0,20 - ed n il numero delle interviste - p. es. 30 - su cui la percentuale è calcolata), la percentuale stessa deve essere rifiutata in quanto il margine dell'errore di campionamento non può essere calcolato con le usuali formule. P. es.: $0,80 \times 0,20 \times 30 = 4,8$. Essendo $4,8 < 9$ la percentuale deve essere rifiutata. Ma se fosse $n=70$, per cui $0,80 \times 0,20 \times 70 = 11,2$, la percentuale potrebbe essere accolta, perché $11,2 > 9$ e il margine di errore potrebbe essere normalmente calcolato con le usuali formule.

ATTENZIONE!

Per le testate a penetrazione bassa (meno di 1-2%) e con lettorato avente connotati di élite culturale o economica, la stima del limite fiduciario va accolta con particolari riserve, data la grande difficoltà di tenere sotto controllo, anche in un campione probabilistico, la varianza.

Si riporta un esempio di tavola fornita per ciascuna edizione dell'indagine "Audicom Print", che contiene gli intervalli fiduciari delle stime di lettura dei settimanali e dei mensili iscritti all'indagine, per il target Carta e/o replica

ESEMPIO Tavola

INTERVALLI FIDUCIARI DELLE STIME DEI SETTIMANALI ISCRITTI ALL'INDAGINE "AUDICOM PRINT"

Lettori Carta e/o Replica

<i>MENSILI</i>	<i>Stime</i>		<i>It. Fid</i>	
	<i>% di penetr.</i>	<i>Stim in '000</i>	<i>% di penetr.</i>	<i>Stim in '000</i>
NOME TESTATA	...	...	...	...
...	...	...	...	...

Universeo Adulti (Nome edizione)	...
Campione (Nome edizione)	...

ESEMPIO Tavola

INTERVALLI FIDUCIARI DELLE STIME DEI MENSILI ISCRITTI ALL'INDAGINE "AUDICOM PRINT"

Lettori Carta e/o Replica

MENSILI	Stime		It. Fid	
	% di penetr.	Stim in '000	% di penetr.	Stim in '000
NOME TESTATA	...	...	...	...
...	...	...	...	...
Universeo Adulti (Nome edizione)	...			
Campione (Nome edizione)	...			

11.4. Attività di cross-recruitment tra Audicom Digital Panel e Audicom Print Survey

Dal 2026, nell'ambito della Ricerca Integrata è prevista un'attività di cross recruitment fra Panel digital e indagine Print.

In particolare, il cross recruitment riguarda:

- la raccolta di lead da Panel digital da invitare all'intervista CAWI dell'indagine Print
- il reclutamento di individui da ricontattare a seguito dell'intervista CAPI Print per aderire al Panel digital

Raccolta di lead da Panel digital per interviste CAWI

Prima dell'avvio di ciascun ciclo di rilevazione, all'interno del campione digitale convenience attivo, viene selezionata una quota di panelisti da invitare alla rilevazione CAWI. La selezione avviene tramite estrazione randomica casuale, con l'obiettivo di assegnare una quota di

individui simili per caratteristiche socio-demografiche e geografiche a ciascuno dei cicli di rilevazione dell'anno.

Ai panelisti digital selezionati viene spedita un'email di invito a partecipare all'indagine CAWI sulla lettura della stampa. L' email di invito contiene una breve descrizione dell'indagine "Audicom Print" e un link univoco (creato ad hoc per ciascun panelista) dal quale accedere al questionario CAWI.

I panelisti digital che accettano di rilasciare l'intervista possono cliccare il collegamento e accedere direttamente al questionario CAWI dell'indagine.

A seguito della trasmissione telematica dei dati delle interviste all'Istituto vengono eseguiti controlli di qualità a tavolino, esaminando criticamente i questionari compilati per controllare la correttezza formale delle registrazioni fatte. Analogamente a tutte le altre interviste dell'indagine, l'eventuale presenza di errori o incoerenze di compilazione viene individuata tramite programmi di cleaning ex-post.

Una volta superati, le interviste vengono incluse nel campione realizzato.

Reclutamenti per Panel digital da interviste CAPI

Il reclutamento di individui da inserire nel Panel digital avviene a seguito della somministrazione dell'intervista CAPI Print, garantendo in questo modo la realizzazione di un campionamento probabilistico su tutto il territorio nazionale.

Prima dell'avvio della rilevazione, viene predisposto lo script da seguire per il reclutamento, con l'obiettivo di raccogliere le ulteriori informazioni (in aggiunta a quelle già raccolte per l'intervista Print) necessarie per essere ricontattati e finalizzare l'adesione.

Lo script consiste in un questionario completamente strutturato, da somministrare al termine dell'intervista. La somministrazione è effettuata utilizzando un PC portatile con il sistema CAPI, analogamente all'intervista dell'indagine "Audicom Print".

Il contatto degli individui da reclutare avviene dagli intervistatori coinvolti nella rilevazione CAPI, pertanto esperti in interviste face to face, che sono adeguatamente formati per questo incarico aggiuntivo.

In particolare, gli intervistatori ricevono un briefing dettagliato sulle attività da svolgere, istruzioni scritte, nonché il materiale da consegnare ai soggetti da reclutare.

Inoltre, vengono comunicate le condizioni di eleggibilità: non sono considerati gli individui che facciano già parte di altri panel e non è possibile reclutare più di una persona per famiglia.

Una volta conclusa l'intervista Print, l'intervistatore spiega in modo molto chiaro le caratteristiche della partecipazione al Panel digital. I soggetti intenzionati forniscono le informazioni per il successivo contatto e procedere con l'adesione.

11.5. Audicom ADV Census - Tracciato record 5.4

Nella tabella sottostante viene riportata la descrizione del tracciato 5.4 per la metrica relativa al contenuto Pubblicitario.

Num. Colonna	Tipo	Codice o Descrizione	Nome	Descrizione e Note
1	Dimensione	Descrizione	Event Date	Data Auditel di visione del contenuto. Formato YYYYMMDD
2	Dimensione	Descrizione	Time Band Start	Inizio intervallo orario di riferimento in formato HHMM (si richiede di avere la mezz'ora, allineata alla giornata Auditel es. 1900).
3	Dimensione	Descrizione	Time Band End	Fine intervallo orario di riferimento in formato HHMM (si richiede di avere la mezz'ora, allineata alla giornata Auditel es. 1930).
4	Dimensione	Descrizione	Publisher	Descrizione dell'editore, proprietario del contenuto erogato mediante la piattaforma. Quando il campo non è popolato (null o stringa vuota) dal broadcaster, dovrà essere usata la stringa "Undefined"
5	Dimensione	Codice	Player Owner Code	Codice fornito da Auditel al proprietario del player e che lo instrumenta (owner). Per il cliente Auditel verrà fornito un Catalogo Statico che assocerà al codice la descrizione del cliente nel tracciato record Auditel.
6	Dimensione	Descrizione	Distribution Platform	Piattaforma che ospita il contenuto. Quando il campo non è popolato (null o stringa vuota) dal broadcaster, dovrà essere usata la stringa "Undefined"
7	Dimensione	Codice	Channel / Property ID	Verrà fornito un Catalogo Statico con la lista dei canali/property e dei relativi editori. Tale catalogo contiene: identificativo e descrizione del canale, concessionaria pubblicitaria di prima chiamata e l'editore (Publisher) proprietario del canale. Questo campo non ha una logica di validazione e deve essere riportato come viene fornito dai broadcaster.
8	Dimensione	Descrizione	Device Class	Verrà fornito un Catalogo Statico. Necessario tracciare i gruppi "Altri Device" e "Non Definito", di seguito altri possibili gruppi (non esaustivi): • PC = 5 • Smartphone = 4 • Tablet = 3 • TV Screen = 2 • Altri Device = 1 • Non Definito = 0
9	Dimensione	Codice	Device OS	Rilevazione del sistema operativo del device. Quando non può essere trovato un valore corretto per qualsiasi ragione, dovrà essere usata la stringa "Undefined"
10	Dimensione	Descrizione	Website / App	Indica il nome del sito o app dove il contenuto è stato visualizzato.
11	Dimensione	Codice	CUSV ID	Codice Univoco Spot Video. Si tratta della possibilità di tracciare ciascun contenuto pubblicitario tramite un codice identificativo univoco, con un catalogo di informazioni aggiuntive sullo spot (progetto CUSV).

Num. Colonna	Tipo	Codice o Descrizione	Nome	Descrizione e Note
12	Dimensione	Codice	Distribution Mode ID	Determina se il video è stato distribuito in modalità Controlled (0) o Uncontrolled (1). Tale distinzione sarà fatta sulla base della baseurl del video confrontata con la lista (whitelist) di siti autorizzati a far visualizzare video dell'editore (in modo owned & operated, embedded o syndication) fornita e continuamente aggiornata dal broadcaster. Verrà fornito un Catalogo Statico.
13	Dimensione	Codice	Out of Country ID	Visualizzazioni effettuate in Italia, in paesi della Comunità Europea, o all'estero. Alle SW House verrà fornito un Catalogo Statico (0 in Italia, 1 all'estero, 2 Comunità Europea). Quando non è possibile avere un valore corretto, dovrà essere usata la stringa "ND".
14	Dimensione	Descrizione	Skippable Duration	Durata in secondi del periodo di tempo per cui l'utente è obbligato a vedere contenuto pubblicitario (-1 significa che il contenuto non è skippabile). Si richiede di valorizzare il campo con una stringa vuota
15	Dimensione	Descrizione	Skipped Flag	Indica se l'utente ha skippato il contenuto pubblicitario (0 se il video non è stato skippato, 1 se il video è stato skippato). Questo campo deve essere valorizzato con una stringa vuota.
16	Dimensione	Descrizione	Content Duration	Durata nominale del contenuto (in secondi). Null, zero, NaN, infinito o altri valori non numerici devono essere sostituiti con -1. Questo valore deve essere riportato con 3 cifre decimali.
17	Dimensione	Descrizione	Primary Adv Agency	Nome della concessionaria pubblicitaria di prima "chiamata". Si richiede di valorizzare il campo con una stringa vuota
18	Dimensione	Descrizione	Adv Content Type	Distinzione se il video è un Pre, Mid o Post Roll. I valori sono Pre-Roll, Post-Roll, Mid-Roll e Undefined (1). Il valore 0 si riferisce a contenuto editoriale.
19	Metrica	Valore numerico	Active Devices No.	Conteggio di Cookie o device id che accedono al contenuto pubblicitario nell'intervallo di tempo (la mezz'ora) e in riferimento ad un oggetto (contenuto/programma/property/canale). Da calcolare sul set di dati delle Stream Views. Questo valore dovrà essere riportato come un numero intero. Il campo non è da popolare, si richiede di inserire una stringa vuota.
20	Metrica	Valore numerico	Stream Start No.	Numero di Video stream avviati. Questo valore deve essere riportato come un numero intero. Il campo è da popolare con il numero di stream avviati, includendo anche quelli invalidi e quelli che non raggiungono i 300 ms.
21	Metrica	Valore numerico	Legitimate Streams No.	Video stream avviati che superano i 300 ms.
22	Metrica	Valore numerico	Legitimate Streams Views 0-25% No.	Video che sono stati visti per una durata compresa tra 0% e il 24,99%.
23	Metrica	Valore numerico	Legitimate Streams Views 25-50% No.	Video che sono stati visti per una durata compresa tra 25% e il 49,99%.
24	Metrica	Valore numerico	Legitimate Streams Views 50-75% No.	Video che sono stati visti per una durata compresa tra 50% e il 74,99%.
25	Metrica	Valore numerico	Legitimate Streams Views 75-100% No.	Video che sono stati visti per una durata compresa 75% e il 100%.

Num. Colonna	Tipo	Codice o Descrizione	Nome	Descrizione e Note
26	Metrica	Valore numerico	Legitimate Streams Views 100% No.	Video che sono stati visti per intero, solo 100%.
27	Metrica	Valore numerico	Total Time Spent (Device)	Tempo totale speso dai device su quel video pubblicitario. Da calcolare sul set di dati delle Legitimate Streams. Questo valore deve essere riportato con 3 cifre decimali.
28	Metrica	Valore numerico	IVT	Video stream avviati (Stream Start) che sono filtrati perché invalidi.
29	Dimensione	Descrizione	Audicom App ID	Indica la classificazione Audicom/Nielsen Market View della Property editoriale. Questo dato è estrapolato dalle informazioni raccolte tramite l'SDK, con l'esclusione delle Connected TV. Per queste ultime, l'App ID sarà ricavata direttamente dalla classificazione Auditel, utilizzando una tabella di conversione.

Campi valorizzati

A seguire vengono riportati i campi previsti ma attualmente non valorizzati sia nel tracciato record prodotto da Auditel che in quello prodotto da Audicom:

- Colonna 12: Distribution Mode ID
- Colonna 14: Skippable Duration
- Colonna 15: Skipped Flag
- Colonna 17: Primary Adv Agency
- Colonna 19: Active Devices No.

A seguire vengono riportati i campi previsti e valorizzati esclusivamente nel tracciato record Audicom:

- Colonna 29: Audicom App ID

A seguire vengono riportati i campi previsti e valorizzati esclusivamente nel tracciato record Auditel:

- Colonna 7: Channel / Property ID