

IBM SPSS Custom Tables 32



Nota

Prima di utilizzare queste informazioni e il prodotto che supportano, leggere le informazioni in [“Informazioni particolari” a pagina 85](#).

Informazioni sul prodotto

Questa edizione si applica alla versione 32, release 0, modifica 0 di IBM® SPSS Statistica e a tutte le release e modifiche successive se non diversamente indicato in nuove edizioni.

© Copyright International Business Machines Corporation .

Indice

Capitolo 1. Tabelle personalizzate.....	1
Interfaccia del Builder delle tabelle.....	1
Interfaccia del Builder delle tabelle.....	1
Tabelle di costruzione.....	1
Tabelle personalizzate: Opzioni Tab.....	15
Tabelle Personalizzate: Scheda Titoli.....	16
Tabelle personalizzate: Test Statistiche Tab.....	17
Tabelle semplici per Variabili categoriali.....	18
Tabelle semplici per Variabili categoriali.....	18
Una variabile singola categoriale.....	19
Tavola di contingenza.....	21
Ordinamento e categorie escluse.....	23
Stack, Nesting e Layers con Variabili categoriali.....	24
Variabili Categoriali categoriali.....	24
Variabili Categoriali Di Nidificazione.....	25
Livelli.....	28
Totali e totali per Variabili categoriali.....	30
Totale semplice per una singola variabile.....	30
Ciò Che Vedete È Quello Che Viene Totalizzato.....	31
Visualizza posizione di Totali.....	31
Totali per le tabelle Nestate.....	32
Totali variabili di strato.....	33
Totali parziali.....	33
Categorie Calcolate per Variabili categoriali.....	35
Categoria semplice Computed.....	35
Nascondere Categorie in una categoria Computed.....	36
Subtotali di riferimento in una categoria Computed.....	37
Utilizzo Categorie Calcolate per visualizzare Subtotali non esaustivi.....	38
Tabelle per Variabili con Categorie condivise.....	39
Tabella dei Conti.....	40
Tabella delle percentuali.....	40
Totali e controllo di categoria.....	41
Nidificazione in tabelle con Categorie condivise.....	42
Statistiche di riepilogo.....	42
Variabile di origine statistiche di riepilogo.....	43
Variabili Stack.....	45
Statistiche totali di riepilogo personalizzate per Variabili categoriali.....	45
Riepilogo Variabili di scala.....	47
Riepilogo Variabili di scala.....	47
Variabili Scale in sovrapposizione.....	48
Statistiche di riepilogo multiple.....	48
Conteggio, N valido e Valori mancanti.....	48
Diversi Sommari per Variabili diverse.....	49
Sommario Gruppo in Categorie.....	50
Intervalli di confidenza.....	52
Statistiche del test.....	53
Statistiche del test.....	53
Test di indipendenza (Chi-quadrato).....	53
Confronto Colonna Significa.....	55
Confronto Proporzioni Colonne.....	58
Una nota sui pesi e gli insiemi a risposta multipla.....	63

Insieme a risposta multipla.....	63
Conteggi, Risposte, Percentuali e Totali.....	64
Utilizzo di Più Set di risposta con Altre Variabili.....	66
Serie di categorie e risposte duplicate multiple.....	67
Significato Testing con Più Set di risposta.....	68
Valori mancanti.....	69
Tabelle senza valori mancanti.....	70
Inclusi i valori mancanti nelle tabelle.....	71
Formattazione e personalizzazione delle tabelle.....	72
Formattazione e personalizzazione delle tabelle.....	72
Formato visualizzazione statistiche di riepilogo.....	72
Visualizza Etichette per statistiche di riepilogo.....	74
Larghezza colonna.....	74
Valore di visualizzazione per celle vuote.....	75
File di esempio.....	76
Informazioni particolari.....	85
Marchi.....	86
Indice analitico.....	89

Capitolo 1. Tabelle personalizzate

Le seguenti funzioni di tabelle personalizzate sono incluse in SPSS Statistiche Standard Edition o l'opzione Tabelle personalizzate.

Interfaccia del Builder delle tabelle

Interfaccia del Builder delle tabelle

Custom Tables utilizza una semplice interfaccia del browser di trascinamento e rilascio che consente di visualizzare in anteprima la propria tabella mentre si selezionano variabili e opzioni. Fornisce inoltre un livello di flessibilità non trovato in una tipica finestra di dialogo, inclusa la possibilità di modificare la dimensione della finestra e la dimensione dei pannelli all'interno della finestra.

Tabelle di costruzione

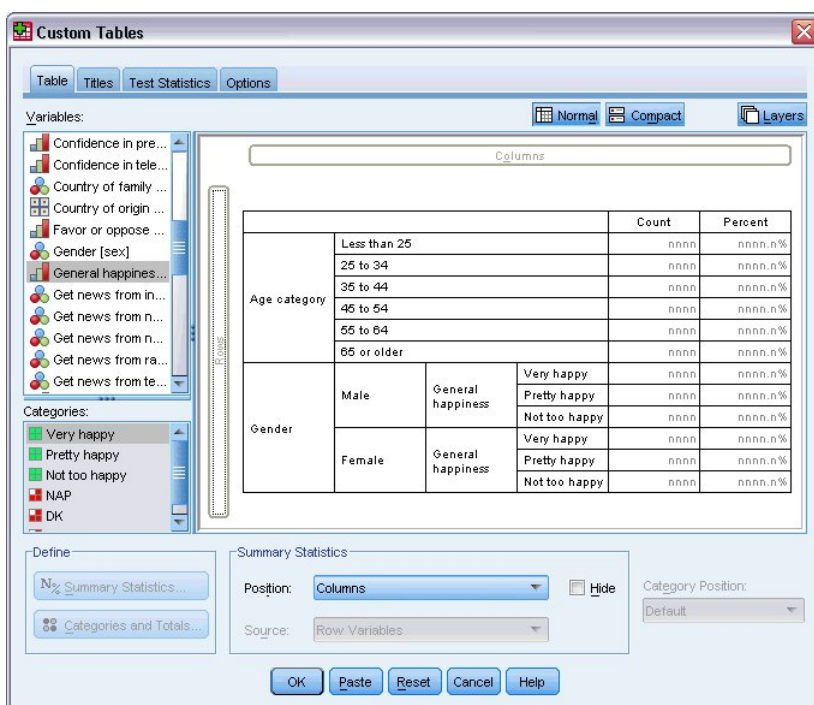


Figura 1. Finestra di dialogo Tabelle personalizzate, tab.

Selezionate le variabili e le misure di riepilogo che appariranno nelle vostre tabelle nella scheda Tabella nel builder da tavolo.

Nota: I campi evidenziati in rosso sono obbligatori. I pulsanti **Incolla** e **OK** vengono abilitati dopo aver inserito valori validi in tutti i campi obbligatori.

Elenco di variabili. Le variabili nel file dei dati vengono visualizzate nel riquadro in alto a sinistra della finestra. Le Tabelle personalizzate distinguono tra due diversi livelli di misurazione per le variabili e le gestisce in modo diverso a seconda del livello di misurazione:

Categoriale. Dati con un numero limitato di categorie o valori distinti (per esempio, genere o religione). Le variabili categoriali possono essere stringa (alfanumeriche) o variabili numeriche che utilizzano codici numerici per rappresentare categorie (ad esempio, 0 = *maschio* e 1 = *femmina*). Vengono denominati anche dati qualitativi. Le variabili categoriali possono essere **nominali** o **ordinale**

- **Nominale.** Una variabile può essere considerata nominale quando i suoi valori rappresentano categorie prive di classificazione intrinseca (ad esempio, il dipartimento della società in cui lavora un dipendente). Degli esempi di variabili nominali includono la regione, il codice postale e l'affiliazione religiosa.
- **Ordinale.** Una variabile può essere trattata come ordinale quando i suoi valori rappresentano categorie con una classificazione intrinseca (ad esempio, livelli di soddisfazione del servizio da molto insoddisfatti a molto soddisfatti). Degli esempi di variabili ordinali includono i punteggi di atteggiamento che rappresentano i gradi di soddisfazione o fiducia e i punteggi di classificazione delle preferenze.

Scala. Dati misurati in base a una scala di intervallo o di rapporto, in cui i valori dei dati indicano sia l'ordine dei valori sia la distanza tra valori. Ad esempio, uno stipendio di \$72.195 è superiore ad uno stipendio di \$52.398 e la distanza tra i due valori è \$19.797. Indicato anche come dati quantitativi o continui.

Le variabili categoriali definiscono le categorie (riga, colonne e strati) nella tabella, e la statistica di riepilogo predefinita è il conteggio (numero di casi in ogni categoria). Ad esempio, una tabella predefinita di una variabile di genere categoriale farebbe semplicemente visualizzare il numero di maschi e il numero di femmine.

Le variabili di scala sono tipicamente riepilogate all'interno di categorie di variabili categoriali e la statistica di riepilogo predefinita è la media. Ad esempio, una tabella di reddito predefinita all'interno delle categorie di genere mostrerebbe il reddito medio per i maschi e il reddito medio per le femmine.

È anche possibile sintetizzare le variabili di scala da sole, senza utilizzare una variabile categoriale per definire i gruppi. Questo è principalmente utile per **impilare** riepiloghi di variabili su più scala. Per ulteriori informazioni, consultare l'argomento [“Variabili Stack”](#) a pagina 4.

Insieme a risposta multipla

Le Tabelle personalizzate supportano anche un tipo speciale di "variabile" chiamata **più serie di risposta**. I set di risposta multipla non sono davvero variabili in senso normale. Non è possibile visualizzarli nell'Editor dei dati e altre procedure non le riconoscono. Gli insiemi a risposta multipla utilizzano più variabili per registrare le risposte alle domande in cui il rispondente può fornire più di una risposta. Tali insiemi vengono considerati come variabili categoriali e nella maggior parte dei casi consentono di eseguire le stesse operazioni previste per questo tipo di variabili. Per ulteriori informazioni, consultare l'argomento [“Insiemi a risposta multipla”](#) a pagina 63.

Un'icona accanto ad ogni variabile nell'elenco delle variabili identifica il tipo di variabile.

Tabella 1. Icone del livello di misurazione				
	Numerico	Stringa	Dati	Ora
Scala (continuo)		n/d		
Ordinale				
Nominale				

Tabella 2. Icone insieme di risposte multiple	
Tipo insieme di risposte multiple	Icona
Insieme di risposte multiple, categorie multiple	
Insieme di risposte multiple, dicotomie multiple	

È possibile modificare il livello di misurazione di una variabile nel builder della tabella facendo clic con il tasto destro del mouse sulla variabile nell'elenco delle variabili e selezionando **Categoriale** o **Scale** dal menu a comparsa. È possibile modificare definitivamente un livello di misurazione di una variabile nella variabile Vista dell'Editor dei dati. Le variabili definite **nominali** o **ordinale** sono trattate come categoriali da tabelle personalizzate.

Categorie. Quando si seleziona una variabile categoriale nell'elenco delle variabili, le categorie definite per la variabile vengono visualizzate nell'elenco Categorie. Queste categorie verranno visualizzate anche sul riquadro canvas quando si utilizza la variabile in una tabella. Se la variabile non ha categorie definite, l'elenco Categorie e il riquadro canvas visualizzeranno due categorie segnaposto: *Categoria 1* e *Categoria 2*.

Le categorie definite visualizzate nel builder table si basano su **etichette di valore**, etichette descrittive assegnate a diversi valori dati (ad esempio valori numerici di 0 e 1, con etichette di valore di *maschio* e *femmina*). È possibile definire etichette di valore in Vista variabile di Data Editor o con Definisci proprietà variabile sul menu Dati nella finestra dell'Editor dei dati.

riquadro Canvas. Si costruiscono una tabella trascinando e rilasciando variabili sulle righe e le colonne del riquadro canvas. Il riquadro canvas visualizza un'anteprima della tabella che verrà creata. Il riquadro canvas non mostra valori dati reali nelle celle, ma dovrebbe fornire una visione abbastanza accurata del layout della tabella finale. Per le variabili categoriali, la tabella effettiva può contenere più categorie rispetto all'anteprima se il file dati contiene valori univocamente per i quali non sono state definite etichette di valore.

- La vista **Normale** visualizza tutte le righe e le colonne che verranno incluse nella tabella, incluse righe e / o colonne per statistiche di riepilogo e categorie di variabili categoriali.
- La vista **Compact** mostra solo le variabili che saranno nella tabella, senza un'anteprima delle righe e delle colonne che la tabella conterrà.

Regole di base e Limitazioni per costruire una Tabella

- Per le variabili categoriali, le statistiche di riepilogo si basano sulla variabile innerpiù nella dimensione di origine delle statistiche.
- La dimensione di origine delle statistiche predefinite (riga o colonna) per le variabili categoriali si basa sull'ordine in cui si trascina e si abbassano le variabili nel riquadro canvas. Ad esempio, se si trascina una variabile al vassoio delle righe prima, la dimensione della riga è la dimensione di origine delle statistiche predefinita.
- Le variabili di scala possono essere riepilogate solo all'interno di categorie della variabile innerpiù nella dimensione della riga o della colonna. (È possibile posizionare la variabile di scala a qualsiasi livello della tabella, ma viene riepilogata a livello innerpiù.)
- Le variabili di scala non possono essere riassunte all'interno di altre variabili di scala. È possibile pilare riepiloghi di variabili di scala multipla o riepilogare le variabili di scala all'interno di categorie di variabili categoriali. Non è possibile nidificare una variabile di scala all'interno di un'altra o mettere una variabile di scala nella dimensione della riga e un'altra variabile di scala nella dimensione della colonna.
- Se una qualsiasi variabile nel dataset attivo contiene più di 12.000 etichette di valore definite, non è possibile utilizzare il builder di tabella per creare tabelle. Se non è necessario includere nelle tabelle variabili che superano tale limite, è possibile definire e applicare insiemi di variabili che escludono tali variabili. Se è necessario includere variabili con più di 12.000 etichette di valori definiti, è possibile utilizzare la CTABLES sintassi del comando per generare le tabelle.

Per Build una tabella

1. Dai menu, scegliere:

Analisi > Tavole > Tabelle personalizzate ...

2. Trascinare e rilasciare una o più variabili alle aree di riga e / o colonne del riquadro canvas.

3. Fare clic su **OK** per creare la tabella.

4. Selezionare (clicca) la variabile sul riquadro canvas.

5. Trascinare la variabile ovunque fuori dal riquadro canvas oppure premere il tasto Elimina.

Per modificare il livello di misurazione di una variabile:

6. Fare clic con il tasto destro del mouse sulla variabile nell'elenco delle variabili (è possibile farlo solo nell'elenco delle variabili, non sui canvas).

7. Selezionare **Categoriale** o **Scale** dal menu a comparsa.

Variabili Stack

Stacking può essere pensato come prendere tavole separate e incollarle insieme nello stesso display. Ad esempio, è possibile visualizzare le informazioni su *Gender* e *Categoria Età* in sezioni separate della stessa tabella.

Alle Variabili di pila

1. Nell'elenco delle variabili, selezionare tutte le variabili che si desidera pilotare, quindi trascinarle e rilasciarle insieme nelle righe o colonne del riquadro canvas.

oppure

2. Trascinare separatamente le variabili di trascinamento, rilasciando ciascuna variabile sopra o sotto le variabili esistenti nelle righe o a destra o a sinistra di variabili esistenti nelle colonne.

Tabella 3. Variabili categoriali sovrapposte		
Variabili	Categorie	Statistica di riepilogo
Variabile 1	Categoria 1	123
	Categoria 2	456
Variabile 2	Categoria 1	123
	Categoria 2	456
	Categoria 3	789

Per ulteriori informazioni, consultare l'argomento [“Variabili Categoriali categoriali”](#) a pagina 24.

Variabili Nidificanti

Nesting, come la crosstabulazione, può mostrare il rapporto tra due variabili categoriali, tranne che una variabile è nidificata all'interno dell'altra nella stessa dimensione. Ad esempio, si potrebbe nidificare *Gender* all'interno di *Categoria di età* nella dimensione della riga, mostrando il numero di maschi e femmine in ogni categoria di età.

È inoltre possibile nidificare una variabile di scala all'interno di una variabile categoriale. Ad esempio, si potrebbe nidificare *Reddito* all'interno di *Gender*, mostrando valori di reddito medi (o mediani o altri misura di riepilogo) per i maschi e le femmine.

Per Nidificare Variabili

1. Trascinare e rilasciare una variabile categoriale nell'area riga o colonna del riquadro canvas.

2. Trascinare e rilasciare una variabile categoriale o di scala a sinistra o a destra della variabile di riga categoriale o superiore o inferiore alla variabile di colonna categoriale.

Tabella 4. Variabili categoriali nidificate		
Variabile 1	Variabile 2	Statistica di riepilogo
Categoria 1	Categoria 1	12
	Categoria 2	34
	Categoria 3	56

Tabella 4. Variabili categoriali nidificate (Continua)		
Variabile 1	Variabile 2	Statistica di riepilogo
Categoria 2	Categoria 1	12
	Categoria 2	34
	Categoria 3	56

Per ulteriori informazioni, consultare l'argomento [“Variabili Categoriali Di Nidificazione”](#) a pagina 25.

Nota: Tavole personalizzate non onorificano l'elaborazione del file suddiviso a strati. Per ottenere lo stesso risultato dei file split stratificati, posizionare le variabili del file suddiviso negli strati di nidificazione più esterni della tabella.

Livelli

È possibile utilizzare gli strati per aggiungere una dimensione di profondità alle proprie tabelle, creando cubi tridimensionali ". Gli strati sono simili a nidificanti o impilabili; la differenza primaria è che solo una categoria di strati è visibile alla volta. Ad esempio, l'utilizzo di *Categoria Età* come variabile di riga e *Gender* come variabile di layer produce una tabella in cui vengono visualizzate informazioni per maschi e femmine in diversi strati della tabella.

Per Creare Layer

1. Clicca su **Layers** nella scheda Tabella nel builder da tavolo per visualizzare l'elenco Layers.
2. Trascinare e rilasciare la variabile o le variabili categoriali che definiranno gli strati nell'elenco Layers.

Non è possibile miscelare le variabili di scala e categoriali nell'elenco Layers. Tutte le variabili devono infatti essere dello stesso tipo. I set di risposta multipla sono trattati come categoriali per l'elenco Layers. Le variabili di scala negli strati sono sempre impilate.

Se si hanno più variabili di layer categoriali, gli strati possono essere impilati o nidificati.

- **Mostra ogni categoria come uno strato** equivale ad impilare. Verrà visualizzato uno strato separato per ogni categoria di ogni variabile di strato. Il numero totale di layer è semplicemente la *somma* del numero di categorie per ogni variabile layer. Ad esempio, se si hanno tre variabili di strato, ognuna con tre categorie, la tabella avrà nove strati.
- **Mostra ogni combinazione di categorie come strato** equivale a strati di nidificazione o crosstabulare. Il numero totale di layer è il *prodotto* del numero di categorie per ogni variabile layer. Ad esempio, se si hanno tre variabili, ognuna con tre categorie, la tabella avrà 27 strati.

Visualizzazione e Nascondere Nomi variabili e / o etichette

Sono disponibili le seguenti opzioni per la visualizzazione di nomi e etichette variabili:

- Mostra solo etichette di variabile. Per qualsiasi variabile senza etichette variabili definite, viene visualizzato il nome variabile. Questa è l'impostazione predefinita.
- Mostra solo nomi variabili.
- Mostrare sia etichette variabili che nomi di variabili.
- Non mostrare nomi variabili o etichette variabili. Sebbene la colonna / riga che contiene l'etichetta o il nome della variabile saranno ancora visualizzate nell'anteprima della tabella sul riquadro canvas, questa colonna / riga non verrà visualizzata nella tabella reale.

Per mostrare o nascondere etichette variabili o nomi variabili:

1. Fare clic con il tasto destro del mouse sulla variabile in anteprima della tabella sul riquadro canvas.
2. Selezionare **Mostra Label variabile** o **Mostra nome variabile** dal menu a comparsa per alternare o disattivare la visualizzazione delle etichette o dei nomi. Un segno di spunta accanto alla selezione indica che verrà visualizzato.

Statistiche di riepilogo

La finestra di dialogo Statistiche di riepilogo consente di:

- Aggiungere e rimuovere statistiche riassuntive da una tabella.
- Cambiare le etichette per le statistiche.
- Modificare l'ordine delle statistiche.
- Modificare il formato delle statistiche, compreso il numero di posizioni decimali.

Le statistiche di riepilogo (e altre opzioni) disponibili qui dipendono dal livello di misurazione della variabile di origine delle statistiche di riepilogo, come visualizzato nella parte superiore della finestra di dialogo. La fonte delle statistiche di riepilogo (la variabile su cui si basano le statistiche di riepilogo) è determinata da:

- **Livello di misurazione.** Se una tabella (o una sezione di tabella in una tabella in sovrapposizione) contiene una variabile di scala, le statistiche di riepilogo si basano sulla variabile di scala.
- **Ordine di selezione variabile.** La dimensione di origine delle statistiche predefinite (riga o colonna) per le variabili categoriali si basa sull'ordine in cui si trascina e si abbassano le variabili sul riquadro canvas. Ad esempio, se si trascina prima una variabile all'area delle righe, la dimensione della riga è la dimensione di origine delle statistiche predefinita.
- **Nidificazione.** Per le variabili categoriali, le statistiche di riepilogo si basano sulla variabile innerpiù nella dimensione di origine delle statistiche.

Una tabella in sovrapposizione può avere più variabili di origine statistiche di riepilogo (sia scala che categoriale), ma ogni sezione di tabella ha una sola fonte di statistiche di riepilogo.

Per modificare la Dimensione di riepilogo Statistiche di riepilogo

1. Selezionare la dimensione (righe, colonne o strati) dall'elenco a discesa **Fonte** nel gruppo di statistiche di riepilogo della scheda Tabella.

Per controllare le statistiche di riepilogo visualizzate in una Tabella

1. Selezionare (clicca) la variabile di origine statistiche di riepilogo sul riquadro canvas della scheda Tabella.
2. Nel gruppo Definisci della scheda Tabella, fare clic su **Statistiche di riepilogo**.
oppure
3. Fare clic con il tasto destro del mouse sulla variabile di origine delle statistiche di riepilogo sul riquadro canvas e selezionare **Statistiche di riepilogo** dal menu a comparsa.
4. Selezionare le statistiche riassuntive che si desidera includere nella tabella. È possibile utilizzare la freccia per spostare le statistiche selezionate dall'elenco Statistiche all'elenco Visualizza oppure si possono trascinare e rilasciare statistiche selezionate dall'elenco Statistiche nell'elenco Visualizza.
5. Fare clic sulle frecce verso l'alto o verso il basso per modificare la posizione di visualizzazione della statistica di riepilogo attualmente selezionata.
6. Selezionare un formato di visualizzazione dall'elenco a discesa Formato per la statistica di riepilogo selezionata.
7. Inserire il numero di decimali da visualizzare nella cella Decimali per la statistica di riepilogo selezionata.
8. Clicca su **Applica alla selezione** per includere le statistiche di riepilogo selezionate per le variabili attualmente selezionate sul riquadro canvas.
9. Clicca su **Applica a tutto** per includere le statistiche di riepilogo selezionate per tutte le variabili impilate dello stesso tipo sul riquadro canvas.

Nota: **Applica a tutto** differisce da **Applica alla selezione** solo per variabili sovrapposte dello stesso tipo già sul riquadro canvas. In entrambi i casi, le statistiche di riepilogo selezionate vengono incluse automaticamente per eventuali variabili sovrapposte dello stesso tipo che si aggiungono alla tabella.

Statistiche riassuntive per le variabili categoriali

Le statistiche di base disponibili per le variabili categoriali sono conteggi e percentuali. È anche possibile specificare statistiche di riepilogo personalizzate per i totali e i totali parziali. Queste statistiche di riepilogo personalizzate includono misure di tendenza centrale (come media e mediana) e dispersione (come la deviazione standard) che possono essere adatte ad alcune variabili categoriali ordinarie. Per ulteriori informazioni, consultare l'argomento [“Statistiche totali di riepilogo personalizzate per Variabili categoriali”](#) a pagina 10.

Conteggio. Numero di casi in ogni cella della tabella o numero di risposte per più set di risposta. Se la ponderazione è in vigore, questo valore è il conteggio ponderato.

- Se la ponderazione è in vigore, il valore è il conteggio ponderato.
- Il conteggio ponderato è lo stesso sia per la ponderazione del dataset globale (**Dati > Casi di peso**) che ponderazione con una variabile di peso base efficace specificata nella scheda **Opzioni** della **Custom Finestra di dialogo tabelle**.

Conteggio non ponderato. Numero non ponderato di casi in ogni cella della tavola. Questo solo differisce dal conteggio se la ponderazione è in vigore.

Conteggio regolato. Il conteggio regolato utilizzato in calcoli di peso base efficaci. Se non si utilizza una variabile peso base efficace (scheda Opzioni), il conteggio regolato è lo stesso del conteggio.

Percentuali di colonna. Percentuali all'interno di ciascuna colonna. Le percentuali in ogni colonna di un sottotetto (per le percentuali semplici) somma al 100%. Le percentuali di colonna sono tipicamente utili solo se si ha una variabile *riga* categoriale.

Percentuali di riga. Percentuali all'interno di ogni riga. Le percentuali in ogni riga di una sottotabella (per le percentuali semplici) somma al 100%. Le percentuali di riga sono tipicamente utili solo se si dispone di una variabile *colonna* categoriale.

Percentuali di riga strato e di colonna strato. Percentuali di righe o colonne (per percentuali semplici) somma al 100% su tutte le sottotavole in un tavolo nidificato. Se la tabella contiene strati, percentuali di righe o colonne somma al 100% su tutte le sottotabelle nidificate in ogni strato.

Percentuali di layer. Percentuali all'interno di ogni strato. Per percentuali semplici, le percentuali di cella all'interno della somma dello strato attualmente visibile al 100%. Se non si hanno variabili di layer, questo equivale a percentuali di tabella.

Percentuali di tabella. Le percentuali per ogni cella sono basate sull'intero tavolo. Tutte le percentuali di cella si basano sullo stesso numero totale di casi e somma al 100% (per le percentuali semplici) su tutta la tabella.

percentuali sottotabella. Le percentuali in ogni cella si basano sul sottotetto. Tutte le percentuali di cella nella sottotabella sono basate lo stesso numero totale di casi e somma al 100% all'interno del sottotetto. Nelle tabelle nidificate, la variabile che precede il livello di nidificazione più intima definisce i sottotavoli. Ad esempio, in una tabella di *stato civile* all'interno di *Gender* all'interno di *Categoria Età*, *Gender* definisce i sottotabelle.

Intervalli di confidenza

- Sono disponibili limiti di confidenza più bassi e superiori per i conteggi, le percentuali, la media, la mediana, i percentili e la somma.
- La stringa di testo "& [Confidence Level]" nell'etichetta include il livello di confidenza nell'etichetta della colonna in tabella.
- L'errore standard è disponibile per i conteggi, le percentuali, la media e la somma.
- Gli intervalli di confidenza e l'errore standard non sono disponibili per più set di risposta.

Livello

Il livello di confidenza per gli intervalli di confidenza, espresso in percentuale. Il valore deve essere maggiore di 0 e minore di 100.

Insiemi a risposta multipla

Più set di risposta possono avere percentuali in base a casi, risposte o conteggi. Per ulteriori informazioni, consultare l'argomento [“Statistiche riassuntive per insiemi a risposta multipla” a pagina 8.](#)

Tabelle impilate

Per i calcoli percentuali, ogni sezione di tabella definita da una variabile di impilamento viene trattata come un tavolo separato. Layer Row, Layer Colonna e percentuali di tabella somma al 100% (per le percentuali semplici) all'interno di ogni sezione di tabella impilata. La base percentuale per diversi calcoli percentuali si basa sui casi in ogni sezione di tabella impilata.

Base percentuale

Le percentuali possono essere calcolate in tre modi diversi, determinati dal trattamento dei valori mancanti nella base computazionale:

Percentuale semplice. Le percentuali si basano sul numero di casi utilizzati nella tabella e si somma sempre al 100%. Se una categoria è esclusa dalla tabella, i casi in quella categoria sono esclusi dalla base. I casi con valori mancanti di sistema sono sempre esclusi dalla base. I casi con valori mancanti degli utenti sono esclusi se le categorie mancanti degli utenti sono escluse dalla tabella (il default) o incluse se le categorie mancanti degli utenti sono incluse nella tabella. Qualsiasi percentuale che non abbia *Valido N* o *Total N* nel suo nome è una percentuale semplice.

% numero totale casi. I casi con i valori mancanti di sistema e quelli mancanti degli utenti vengono aggiunti alla base percentuale semplice. È possibile che la somma delle percentuali sia inferiore al 100%.

% casi validi. I casi con valori mancanti dell'utente vengono rimossi dalla base percentuale semplice anche se le categorie mancanti degli utenti sono incluse nella tabella.

Nota: i casi in categorie escluse manualmente diverse dalle categorie mancanti degli utenti sono sempre esclusi dalla base.

Statistiche riassuntive per insiemi a risposta multipla

Le seguenti statistiche di riepilogo aggiuntive sono disponibili per più set di risposta.

Col / Riga / Risposte Layer%. Percentuale basata sulle risposte.

Col / Riga / Risposte Layer% (Base: Conteggio). Le risposte sono il numeratore e il conteggio totale è il denominatore.

Col / Riga / Layer Count% (Base: Risposte). Il conteggio è il numeratore e le risposte totali sono il denominatore.

Layer Col / Row Responses%. Percentuale su sottotabelle. Percentuale basata sulle risposte.

Layer Col / Riga Risposte% (Base: Conteggio). Percentuali su sottotavoli. Le risposte sono il numeratore e il conteggio totale è il denominatore.

Strato Col/RowResponses % (Base: risposte). Percentuali su sottotavoli. Il conteggio è il numeratore e le risposte totali sono il denominatore.

Risposte. Conteggio delle risposte.

Sottotavolo / Tabella Risposte%. Percentuale basata sulle risposte.

Sottotavolo / Tabella Risposte% (Base: Conteggio). Le risposte sono il numeratore e il conteggio totale è il denominatore.

Sottotetto / Conteggio tabelle% (Base: Risposte). Il conteggio è il numeratore e le risposte totali sono il denominatore.

Statistiche di riepilogo per variabili di scala e Totali personalizzati categoriali

Oltre ai conteggi e alle percentuali disponibili per le variabili categoriali, sono disponibili le seguenti statistiche di riepilogo per le variabili di scala e come riepiloghi totali e subtotali personalizzati per variabili categoriali. Queste statistiche di riepilogo non sono disponibili per più serie di risposta o variabili stringa (alfanumeriche).

Media. Media aritmetica; la somma divisa per il numero di casi.

Mediana. Valore al di sopra e al di sotto del quale ricade la metà dei casi; il 50° percentile.

Modalità. Valore più ricorrente. Se c'è una cravatta, viene mostrato il valore più piccolo.

Minimo. Valore più piccolo (più basso).

Massimo. Valore più grande (più alto).

Mancante. Conteggio dei valori mancanti (sia di utente che di sistema).

Percentile. È possibile includere il 5°, il 25°, il 75°, il 95° e/o il 99° percentile.

Intervallo. Differenza tra valori massimi e minimi.

deviazione standard. Misura della dispersione intorno alla media. In una distribuzione normale, il 68% dei casi rientra all'interno di una deviazione standard della media e il 95% dei casi rientra nell'ambito di due deviazioni standard. Ad esempio, se l'età media è di 45 anni, con una deviazione standard del 10, il 95% dei casi sarebbe compreso tra 25 e 65 in una distribuzione normale (radice quadrata della varianza).

Somma. Somma dei valori.

Somma percentuale. Percentuali in base alle somme. Disponibile per righe e colonne (all'interno di sottotabelle), intere righe e colonne (attraverso sottotabelle), strati, sottotabelle e tabelle intere.

Total N. Conteggio dei valori non mancanti, mancanti di utenti e di sistema. Non include i casi nelle categorie manualmente escluse diverse dalle categorie mancanti degli utenti.

Totale rettificato N. Il totale corretto N utilizzato in calcoli di peso base efficace. Se non si utilizza una variabile peso base efficace (scheda Opzioni), il totale corretto N è lo stesso del totale N. Questa statistica non è disponibile per più set di risposta.

N. valido N. Conteggio dei valori non mancanti. Non include i casi nelle categorie manualmente escluse diverse dalle categorie mancanti degli utenti.

Aggiustato Valido N. Il N corretto aggiustato utilizzato in calcoli di peso base efficaci. Se non si utilizza una variabile di peso base efficace (scheda Opzioni), il N corretto aggiustato è uguale a quello valido N. Questa statistica non è disponibile per più set di risposta.

Varianza. Una misura di dispersione intorno alla media, pari alla somma delle deviazioni quadrate dalla media divisa per uno inferiore al numero di casi. La varianza è misurata in unità che sono la piazza di quelle della variabile stessa (il quadrato della deviazione standard).

Intervalli di confidenza

- Sono disponibili limiti di confidenza più bassi e superiori per i conteggi, le percentuali, la media, la mediana, i percentili e la somma.
- La stringa di testo "& [Confidence Level]" nell'etichetta include il livello di confidenza nell'etichetta della colonna in tabella.
- L'errore standard è disponibile per i conteggi, le percentuali, la media e la somma.
- Gli intervalli di confidenza e l'errore standard non sono disponibili per più set di risposta.

Livello

Il livello di confidenza per gli intervalli di confidenza, espresso in percentuale. Il valore deve essere maggiore di 0 e minore di 100.

Tabelle impilate

Ogni sezione di tabella definita da una variabile di impilamento viene trattata come una tabella separata e le statistiche di riepilogo vengono calcolate di conseguenza.

Statistiche totali di riepilogo personalizzate per Variabili categoriali

Per tabelle di variabili categoriali che contengono totali o parziali, è possibile avere statistiche di riepilogo diverse rispetto ai riepiloghi visualizzati per ogni categoria. Ad esempio, si potrebbero visualizzare conteggi e percentuali di colonne per una variabile riga categoriale ordinale e visualizzare la mediana per la statistica "totale".

Per creare una tabella per una variabile categoriale con una statistica di riepilogo totale personalizzata:

1. Dai menu, scegliere:

Analisi > Tavole > Tabelle personalizzate ...

Il costruttore di tabelle si aprirà.

2. Trascinare e rilasciare una variabile categoriale nell'area Rows o Columns delle canvas.
3. Fare clic con il tasto destro del mouse sulla variabile sui canvas e selezionare **Categorie e Totali** dal menu a comparsa.
4. Fare clic (controllare) la casella di controllo **Totale**, quindi fare clic su **Applica**.
5. Fare clic con il tasto destro del mouse sulla variabile di nuovo sui canvas e selezionare **Statistiche di riepilogo** dal menu a comparsa.
6. Clicca (spunta) **Statistiche riepilogo personalizzate per Totali e Subtotali**, quindi selezionare le statistiche di riepilogo personalizzate desiderate.

Per impostazione predefinita, tutte le statistiche di riepilogo, inclusi i riepiloghi personalizzati, vengono visualizzate nella dimensione opposta dalla dimensione contenente la variabile categoriale. Ad esempio, se si ha una variabile riga categoriale, le statistiche di riepilogo definiscono colonne nella tabella, come in:

<i>Tabella 5. Categorie variabili ordinali in righe, conteggio delle statistiche di riepilogo e media nelle colonne</i>			
Variabili	Categorie	Conteggio	Media
Variabile 1	1 Concorsi	196	2.29
	2 Neutro	936	
	3 Disaccordo	744	
	Totale	1876	

Per visualizzare le statistiche di riepilogo nella stessa dimensione della variabile categoriale:

7. Nella scheda Tabella nel builder table, nel gruppo Riepilogo statistiche, selezionare la dimensione dall'elenco a discesa Posizione.

Ad esempio, se la variabile categoriale viene visualizzata nelle righe, selezionare **Righe** dall'elenco a discesa.

Formati di visualizzazione delle statistiche riassuntive

Sono disponibili le seguenti opzioni di formato di visualizzazione:

nnnn. Numerico semplice.

nnnn%. Segno percentuale accostato alla fine del valore.

Automatico. Formato di visualizzazione variabile, compreso il numero di decimali.

N=nnnn. Visualizza N= prima del valore. Questo può essere utile per i conteggi, validi Ne il totale N nelle tabelle in cui le etichette statistiche di riepilogo non vengono visualizzate.

(nnnn). Tutti i valori racchiusi tra parentesi.

(nnnn) (neg. valore). Solo valori negativi racchiusi tra parentesi.

(nnnn%). Tutti i valori racchiusi tra parentesi e un segno percentuale accostato alla fine dei valori.

n, nnn.n. Formato virgola. Virgola utilizzata come separatore di raggruppamento e periodo utilizzato come indicatore decimale indipendentemente dalle impostazioni locali.

n.nnn, n. Formato dot. Periodo utilizzato come separatore di raggruppamento e virgola utilizzato come indicatore decimale indipendentemente dalle impostazioni locali.

\$n, nnn.n. Formato del dollaro. Segno del dollaro visualizzato davanti al valore; virgola utilizzata come separatore di raggruppamento e periodo utilizzato come indicatore decimale indipendentemente dalle impostazioni locali.

CCA, CCB, CCC, CCD, CCE. Formati di valuta personalizzati. L'attuale formato definito per ogni valuta personalizzata viene visualizzato nell'elenco. Questi formati sono definiti sulla scheda Currency nella finestra di dialogo Opzioni (Modifica menu, Opzioni).

Regole generali e Limitazioni

- Ad eccezione di Auto, il numero di decimali è determinato dall'impostazione della colonna Decimals.
- Ad eccezione dei formati da virgola, dollaro e dot, l'indicatore decimale utilizzato è quello definito per il locale corrente nel proprio pannello di controllo delle Opzioni Regionali di Windows.
- Sebbene comma/dollaro e dot visualizzeranno rispettivamente una virgola o un periodo come separatore di raggruppamento, non esiste un formato di visualizzazione disponibile al momento della creazione per visualizzare un separatore di raggruppamento in base alle impostazioni del locale corrente (definite nel pannello di controllo delle Opzioni regionali di Windows).

Categorie e totali

La Finestra di dialogo Categorie e Totali consente di:

- Riordinare ed escludere le categorie.
- Inserisci subtotali e totali.
- Inserire categorie calcolate.
- Includere o escludere le categorie vuote.
- Includere o escludere categorie definite come contenenti valori mancanti.
- Includere o escludere categorie che non hanno etichette di valore definite.
- Questa finestra di dialogo è disponibile solo per variabili categoriali e set di risposta multipla. e non è disponibile per le variabili di scala.
- Per più variabili selezionate con diverse categorie, non è possibile inserire sottototali, inserire categorie calcolate, escludere categorie o categorie di riordino manualmente. Questo si verifica solo se si selezionano più variabili nell'anteprima canvas e si accede a questa finestra di dialogo per tutte le variabili selezionate contemporaneamente. È comunque possibile eseguire queste azioni per ciascuna variabile separatamente.
- Per le variabili senza etichette di valore definite è possibile ordinare solo categorie e inserire totali.

Per accedere alle Categorie e Totali Dialog Box

1. Trascinare e rilasciare una variabile categoriale o più serie di risposta sul riquadro canvas.
2. Fare clic con il tasto destro del mouse sulla variabile sul riquadro canvas e selezionare **Categorie e Totali** dal menu a comparsa.

oppure

3. Selezionare (cliccare) la variabile sul riquadro canvas, quindi fare clic su **Categorie e totali** nel gruppo Definisci nella scheda Tabella.

È inoltre possibile selezionare più variabili categoriali nella stessa dimensione sul riquadro canvas:

4. Fare clic, tenendo premuto Ctrl, su ciascuna variabile nel riquadro dell'area.

oppure

5. Clicca all'esterno dell'anteprima della tabella sul riquadro canvas, quindi clicca e trascina per selezionare l'area che include le variabili che si desidera selezionare.

oppure

6. Fare clic con il tasto destro del mouse su qualsiasi variabile in una dimensione e selezionare **Seleziona tutto [dimensione] Variabili** per selezionare tutte le variabili in tale dimensione.

Per Riordinare Categorie

Per riordinare manualmente le categorie:

1. Selezionare una categoria dall'elenco.
 2. Clicca sulla freccia in alto o in basso per spostare la categoria in alto o in basso nell'elenco.
- oppure
3. Clicca nella colonna Valore (s) per la categoria, e trascina e rilascia in una posizione diversa.

Per Escludere Categorie

1. Selezionare una categoria dall'elenco.
 2. Clicca sulla freccia accanto all'elenco Escludi.
- oppure
3. Clicca nella colonna Valore (s) per la categoria e trascina e rilascia ovunque fuori dall'elenco.

Se si escluda qualsiasi categorie, saranno escluse anche le categorie senza etichette di valore definite.

Per Ordinare Categorie

È possibile ordinare categorie per valore dati, etichetta del valore, conteggio delle celle o statistica di riepilogo in ordine crescente o decrescente.

1. Nel gruppo Categorie di ordinamento, clicca sull'elenco a discesa Per selezionare il criterio di ordinamento che si desidera utilizzare: valore, etichetta, conteggio o statistica di riepilogo (come media, mediana o modalità). Le statistiche di riepilogo disponibili per l'ordinamento dipendono dalle statistiche di riepilogo selezionate per visualizzare in tabella.
2. Clicca sull'elenco a discesa dell'ordine per selezionare l'ordine di ordinamento (ascendente o discendente).

Le categorie di ordinamento non sono disponibili se si hanno escluso qualsiasi categorie.

Totali parziali

1. Selezionare (clicca) la categoria nell'elenco che è l'ultima categoria nella gamma di categorie che si desidera includere nel sottototale.
2. Clicca su **Aggiungi Sottotale ...**
3. Nella finestra di dialogo Definisci Sottotale, opzionalmente modificare il testo dell'etichetta sottotale.
4. Per mostrare solo un sottotale e sopprimere la visualizzazione delle categorie che definiscono il sottotale, selezionare **Nascondi categorie sottototali dalla tabella**.
5. Clicca su **Continua** per aggiungere il sottotale.

Totali

1. Fare clic sulla casella di controllo **Totale** . È anche possibile modificare il testo dell'etichetta totale.

Se la variabile selezionata viene nidificata all'interno di un'altra variabile, i totali verranno inseriti per ogni sottototale.

Visualizza posizione per Totali e Subtotali

I totali e i totali parziali possono essere visualizzati sopra o sotto le categorie incluse in ogni totale.

- Se **Sotto** viene selezionato nel gruppo Totals e Subtotals Aspetto, i totali appaiono sopra ogni sottotetto e tutte le categorie sopra e compresa la categoria selezionata (ma al di sotto di eventuali subtotali precedenti) sono incluse in ogni subtotale.
- Se **Sopra** viene selezionato nel gruppo Totals and Subtotals Aspetto, i totali appaiono al di sotto di ogni sottotabella e tutte le categorie sottostanti e compresa la categoria selezionata (ma soprattutto i subtotali precedenti) sono incluse in ogni subtotale.

Importante: è necessario selezionare la posizione di visualizzazione per i subtotali prima di definire eventuali subtotali. La modifica della posizione del display interessa tutti i subtotali (non solo il subtotale attualmente selezionato) e inoltre *cambia le categorie incluse nei subtotali*.

Categorie Calcolate

È possibile visualizzare le categorie calcolate da statistiche di riepilogo, totali, totali e / o costanti. Per ulteriori informazioni, consultare l'argomento [“Categorie Calcolate” a pagina 13](#).

Statistiche di riepilogo Total e Subtotali

È possibile visualizzare statistiche diverse da "totali" nelle aree Totali e Subtotali della tabella utilizzando la finestra di dialogo Statistiche di riepilogo. Per ulteriori informazioni, consultare l'argomento [“Statistiche riassuntive per le variabili categoriali” a pagina 7](#).

Nota: se si selezionano più statistiche totali personalizzate che si trovano anche nel corpo della tabella e si nasconde le etichette statistiche, allora i totali vengono riordinati nello stesso ordine del corpo della tabella - e visto che le etichette non vengono visualizzate, non si può sapere cosa rappresenta in realtà ogni statistica totale. In generale, selezionando più statistiche e nascondendo le etichette statistiche probabilmente non è una buona idea.

Totali, Subtotali e Categorie escluse

I casi da categorie escluse non sono inclusi nel calcolo dei totali.

Valori Mancanti, Categorie vuote e Valori senza Value Labels

Valori mancanti. Questo controlla la visualizzazione dei valori di **user - missing** o valori definiti come contenenti valori mancanti (ad esempio, un codice di 99 da rappresentare "non applicabile" per la gravidanza nei maschi). I valori mancanti definiti dall'utente vengono esclusi per impostazione predefinita. Selezionare (controllare) questa opzione per includere le categorie mancanti degli utenti nelle tabelle. Sebbene la variabile possa contenere più di una categoria di valore mancante, l'anteprima della tabella sui canvas visualizzerà solo una categoria di valore mancante generica. Tutte le categorie mancanti definite dall'utente saranno incluse nella tabella. **I valori mancanti di sistema** (celle vuote per variabili numeriche nell'editor dei dati) sono sempre esclusi.

Categorie vuote. Le categorie vuote sono categorie con etichette di valore definite ma nessun caso in quella categoria per una determinata tabella o sottotabella. Per impostazione predefinita, le categorie vuote sono incluse nelle tabelle. Deselezionare questa opzione per escludere le categorie mancanti dalla tabella.

Altri valori trovati quando vengono scansionati i dati. Per impostazione predefinita, i valori di categoria nel file dati che non hanno etichette di valore definite sono automaticamente inclusi nelle tabelle. Deselezionare questa opzione per escludere i valori senza etichette di valore definite dalla tabella. Se si escluda qualsiasi categoria con etichette di valore definite, sono escluse anche le categorie senza etichette di valore definite.

Categorie Calcolate

Oltre a visualizzare i risultati aggregati delle statistiche di riepilogo, una tabella può visualizzare una o più categorie calcolate da questi risultati aggregati, da valori costanti, da subtotali e totali, o da una combinazione di essi. I risultati sono noti come categorie computate o postcompute. Una categoria computata agisce come una categoria in un'unica variabile con le seguenti similitudini e differenze:

- Una categoria calcolata è posizionata come le altre categorie.
- Una categoria calcolata opera sulle stesse statistiche delle altre categorie.
- Le categorie calcolate non influiscono su subtotali, totali o test di significatività.
- Per impostazione predefinita, i valori delle categorie calcolate utilizzano la stessa formattazione per le statistiche di riepilogo come le altre categorie. È possibile sovrascrivere il formato quando si definisce la categoria computata.

Poiché le categorie calcolate possono essere utilizzate a risultati aggregati totali, possono essere simili a subtotali. Tuttavia, le categorie calcolate hanno i seguenti vantaggi rispetto ai subtotali:

- Le categorie calcolate possono essere calcolate dai risultati di altri subtotali.
- Le categorie calcolate possono sovrapporsi tra loro, operando sulle stesse (o alcune delle stesse) categorie.
- Le categorie calcolate non devono includere valori da tutte le altre categorie sopra o sotto la categoria computata. Cioè le categorie calcolate non sono esaustive.
- Le categorie calcolate possono includere valori da categorie non adiacenti.

A differenza dei totali e dei subtotali, le categorie calcolate sono calcolate a partire dai dati aggregati piuttosto che dai dati originali. Pertanto, i valori delle categorie calcolate possono non corrispondere ai risultati dei totali e dei subtotali. Inoltre, poiché si ha l'opzione di nascondere le categorie di origine quando si definisce la categoria computata, potrebbe essere difficile interpretare i subtotali nella tabella risultante. Se si utilizzano categorie calcolate, si consiglia di specificare etichette personalizzate per i subtotali.

Per definire una categoria Computed

Le categorie calcolate vengono aggiunte dalla finestra di dialogo Categorie e Totali. Per informazioni sull'accesso a quella finestra di dialogo, consultare l'argomento [“Categorie e totali”](#) a pagina 11 .

1. Nella finestra di dialogo Categorie e Totali, fare clic su **Aggiungi categoria ...**
2. In **Label for Computed Category**, specificare un'etichetta per la categoria computata. È possibile trascinare le categorie dall'elenco Categorie per includere le etichette per quelle categorie.
3. Costruire un'espressione selezionando categorie e / o totali e subtotali e utilizzando gli operatori per definire le categorie calcolate. È anche possibile immettere valori costanti (ad esempio, 500) da includere nell'espressione.
4. Per mostrare solo una categoria computata e sopprimere la visualizzazione delle categorie che definiscono la categoria del calcolo, selezionare **Categorie Nascondi utilizzate in espressione dalla tabella**.
5. Fare clic sulla scheda **Visualizza formati Formati** per modificare il formato di visualizzazione e il numero di posizioni decimali per la categoria computata. Per ulteriori informazioni, consultare l'argomento [“Visualizza Formati per Categorie Compute”](#) a pagina 14.
6. Clicca su **Continua** per aggiungere la categoria computata.

Visualizza Formati per Categorie Compute

Per impostazione predefinita, una categoria computata utilizza lo stesso formato di visualizzazione e il numero di posizioni decimali come le altre categorie nella variabile. È possibile sovrascrivere questi nella scheda Visualizza Formati nella finestra di dialogo Categoria Computed. La Scheda Formati di visualizzazione elenca le statistiche di riepilogo correnti su cui opera la categoria computata in aggiunta ai formati di visualizzazione e al numero di posizioni decimali per quelle statistiche.

Per ogni statistica di riepilogo è possibile:

1. Selezionare un formato di visualizzazione dall'elenco a discesa Formato per la statistica di riepilogo. Per un elenco completo dei formati di visualizzazione, consultare l'argomento [“Formati di visualizzazione delle statistiche riassuntive”](#) a pagina 10 .
2. Inserire il numero di decimali da visualizzare nella cella Decimali per la statistica di riepilogo selezionata.

Tabelle di Variabili con Categorie condivise (Tavole Comperimetriche)

Le indagini spesso contengono molte domande con un insieme comune di possibili risposte. È possibile utilizzare la sovrapposizione per visualizzare queste variabili correlate nella stessa tabella ed è possibile visualizzare le categorie di risposta condivise nelle colonne della tabella.

Per creare una Tabella per Variabili multiple con Categorie Condivise

Tabella 6. Variabili impilate con categorie di risposta condivise nelle colonne			
Variabili	Categoria 1	Categoria 2	Categoria 3
Variabile 1	12	34	56
Variabile 2	56	12	34
Variabile 3	34	56	12

Per ulteriori informazioni, consultare l'argomento [“Tabelle per Variabili con Categorie condivise”](#) a pagina 39.

Personalizzazione del Builder delle tabelle

A differenza delle finestre di dialogo standard, è possibile modificare la dimensione del builder di tabella nello stesso modo in cui è possibile modificare le dimensioni di qualsiasi finestra standard:

1. Clicca e trascina la parte superiore, inferiore, laterale o qualsiasi angolo del costruttore di tavolo per diminuire o aumentare la sua dimensione.
Nella scheda Tabella è possibile modificare anche la dimensione dell'elenco delle variabili, l'elenco Categorie e il riquadro canvas.
2. Clicca e trascina la barra orizzontale tra l'elenco delle variabili e l'elenco Categorie per rendere le liste più lunghe o più brevi. Spostarlo rende più lungo l'elenco delle variabili e la lista Categorie più breve. Spostarlo fa il contrario.
3. Clicca e trascina la barra verticale tra l'elenco delle variabili e l'elenco Categorie dal riquadro canvas per rendere le liste più ampie o più strette. La canvas si ridimensiona automaticamente per adattarsi allo spazio rimanente.

Tabelle personalizzate: Opzioni Tab

Aspetto cella dati

Controlla ciò che viene visualizzato in celle e celle vuote per le quali le statistiche non possono essere calcolate.

Celle vuote

Per le celle da tavolo che non contengono casi (conteggio cellulare pari a 0) è possibile selezionare una delle tre opzioni di visualizzazione: zero, blank o un valore di testo specificato. Il valore di testo può essere lungo fino a 255 caratteri.

Impossibile calcolare le statistiche

Testo che viene visualizzato se una statistica non può essere calcolata (ad esempio, la media per una categoria senza casi). Il valore di testo può essere lungo fino a 255 caratteri. Il valore predefinito è un periodo (.).

Larghezza delle colonne dati

Controlla la larghezza minima e massima della colonna per le colonne dei dati. Questa impostazione non influisce sulle larghezze delle colonne per le etichette di riga.

Impostazioni TableLook

Utilizza la specifica della larghezza della colonna di dati dal TableLook TableLook. È possibile creare il proprio TableLook predefinito personalizzato da utilizzare quando vengono create nuove tabelle ed è possibile controllare la larghezza delle colonne delle etichette di riga e delle colonne di dati con un TableLook.

Personalizzato

Sovrascrive le impostazioni TableLook predefinite per la larghezza della colonna di dati.
Specificare le larghezze di colonna dei dati minimi e massimi per la tabella e l'unità di misurazione: punti, pollici o centimetri.

Valori mancanti per le variabili di scala

Per le tabelle con due o più variabili di scala, controlla la gestione dei dati mancanti per le statistiche variabili su scala.

Massimizza l'uso dei dati disponibili (eliminazione variabile per variabile)

Tutti i casi con valori validi per ogni variabile di scala sono inclusi nelle statistiche di riepilogo per quella variabile.

Usa base dei casi congruente nelle variabili di scala (eliminazione listwise)

I casi con valori mancanti per eventuali variabili di scala nella tabella sono esclusi dalle statistiche di riepilogo per tutte le variabili di scala presenti nella tabella.

Base effettiva

Se si ha una variabile che rappresenta i pesi di regolazione piuttosto che i pesi di frequenza, è possibile utilizzare quella variabile come una variabile peso base efficace. Il concetto di base efficace o di efficacia del campione efficace è basato sui metodi per l'analisi dei dati provenienti da campioni complessi. Un peso base efficace consente una gestione approssimativa dell'inferenza statistica in analisi che comporta aggiustamenti ad hoc dei dati provenienti da semplici disegni di campionamento casuali utilizzando i pesi di regolazione.

- Il peso base efficace influenza i valori di statistica di riepilogo ponderati e i mezzi di colonna e le proporzioni delle proporzioni della colonna.
- Se la ponderazione viene attivata per il dataset, la variabile peso del dataset viene ignorata e i risultati sono ponderati dalla variabile peso base efficace.
- La variabile peso base efficace deve essere numerica.
- I casi con valori di peso negativi, un valore di peso di 0 o valori di peso mancanti sono esclusi da tutti i risultati.

Conta risposte duplicate per gli insiemi a categorie multiple

Una risposta duplicata è la stessa risposta per due o più variabili nel set di più categorie. Per impostazione predefinita, le risposte duplicate non sono conteggiate.

Nascondi conteggi piccoli

È possibile scegliere di nascondere i conteggi che sono inferiori a un numero intero specificato.

- I valori nascosti vengono visualizzati come **< N**, dove **N** è il numero intero specificato.
- Il numero intero specificato deve essere superiore o uguale a 2.
- Se la ponderazione viene attivata per il dataset o viene specificata una variabile peso base efficace, viene utilizzato il valore ponderato.

Pesi e Rounding

- Se si utilizza **Dati > Casi di peso** ai casi di peso, i pesi non interi sono arrotondati a livello di cella o di categoria per i test di significatività, gli intervalli di confidenza e gli errori standard.
- Se si seleziona **Utilizza variabile peso base effettiva**, i valori di peso non interi nella variabile peso selezionata non sono arrotondati.
- Se entrambi sono specificati, viene utilizzata la variabile peso base efficace e i pesi non interi non sono arrotondati.

Tabelle Personalizzate: Scheda Titoli

La Scheda Titoli controlla la visualizzazione di titoli, didascalie e etichette d'angolo.

Titolo. Testo che viene visualizzato sopra la tabella.

Didascalia. Testo che viene visualizzato sotto la tabella e sopra eventuali note.

Angolo. Testo che viene visualizzato nell'angolo in alto a sinistra della tabella. Il testo dell'angolo viene visualizzato solo se la tabella contiene variabili di riga e se la proprietà etichetta della quota di riga della tabella pivot è impostata su **Nestati**. Questa *non* è l'impostazione predefinita TableLook .

È possibile includere i seguenti valori generati automaticamente nel titolo della tabella, didascalia o etichetta d'angolo:

Data. Anno corrente, mese e giorno visualizzati in un formato basato sulle tue attuali impostazioni di Opzioni regionali di Windows.

Ora. Ora corrente, minuto e seconda visualizzata in un formato basato sulle tue attuali impostazioni di Opzioni regionali di Windows.

Espressione tabella. Variabili utilizzate nella tabella e come vengono utilizzate nella tabella. Se una variabile ha un'etichetta variabile definita, viene visualizzata l'etichetta. Nella tabella generata, i seguenti simboli indicano come vengono utilizzate le variabili nella tabella:

- + indica variabili sovrapposte.
- > indica la nidificazione.
- **BY** indica la crosstabulazione o gli strati.

Tabelle personalizzate: Test Statistiche Tab

La scheda Statistiche del test fornisce test di significatività per tabelle personalizzate.

Questi test non sono disponibili per tabelle in cui le etichette di categoria vengono spostate fuori dalla loro dimensione di tabella predefinita o per categorie calcolate.

Colonna Medie e Colonna Proporzioni Test

I test delle colonne sono disponibili per le variabili di scala. I test delle proporzioni della colonna sono disponibili per le variabili categoriali.

Confronta medie di colonna

Test pairwise dell'eguaglianza dei mezzi di colonna. La tabella deve avere una variabile categoriale nelle colonne e una variabile di scala come il livello più alto delle righe. La tabella deve includere la media come statistica di riepilogo.

Per le variabili categoriali ordinarie, la varianza può essere stimata da tutte le categorie o da solo le categorie confrontabili. Per le variabili di risposta multiple, la varianza per il test dei mezzi è sempre basata sulle sole categorie che vengono confrontati.

Confronta proporzioni di colonna

Test pairwise dell'eguaglianza delle proporzioni delle colonne. La tabella deve avere almeno una variabile categoriale sia nelle colonne che nelle righe. La tabella deve includere conteggi o percentuali di colonne.

Identifica differenze significative

Per i test di proporzioni colonne e colonne, è possibile visualizzare risultati significativi in una tabella separata o nella tabella principale.

In una tabella separata

I risultati dei test di significatività vengono visualizzati in una tabella separata. Se due valori sono significativamente diversi, la cella corrispondente al valore più grande visualizza una chiave che identifica la colonna con il valore più piccolo.

Visualizza valori di significatività

I valori di significatività vengono visualizzati tra parentesi dopo ogni valore chiave nella cella. Questa opzione è disponibile solo quando i risultati significativi vengono visualizzati in una tabella separata.

Nella tabella principale

I risultati del test di significatività sono visualizzati nella tabella principale. Ogni categoria di colonna nella tabella è identificata con una chiave alfabetica. Per ogni coppia significativa, la

chiave della categoria con la media di colonna più piccola o proporzione compare nella categoria con la colonna più grande o proporzione.

- Quando si supera una chiave nella cella etichetta della colonna in una tabella pivot, vengono evidenziate tutte le celle della tabella con tale tasto di significazione. Per una tabella con più variabili nella dimensione della colonna, vengono evidenziate solo le celle in quella sottotabella.
- Per selezionare tutte le celle in una tabella (o sottotabella) che hanno la stessa chiave di significato, fare clic con il tasto destro del mouse sulla cella dell'etichetta della colonna e scegliere **Seleziona > Selezionare tutte le celle con questa chiave di significato**.

Utilizza indici stile APA

Identificare differenze significative con la formattazione in stile APA che utilizza le lettere di sottoscrittura. Se due valori sono significativamente diversi, tali valori visualizzano diverse lettere di sottoscrittura. Questi sottoscritti non sono note. Quando questa opzione è attiva, lo stile di piè di pagina definito nel TableLook corrente viene sovrascritto e le note a piè di pagina vengono visualizzate come numeri in apice. Per selezionare tutte le celle della stessa riga con la stessa chiave di significato, clicca con il tasto destro del mouse su una cella che ha una chiave di significato e scegliere **Seleziona celle con significato simile**

Livelli di significatività

Il livello di significatività per i mezzi di colonna e le proporzioni delle colonne.

- Il valore deve essere superiore a 0 e inferiore a 1.
- Se si specificano due livelli di significati, le lettere maiuscoli vengono utilizzate per identificare valori di significatività inferiori o uguali al livello più piccolo. Le lettere minuscoli sono utilizzate per identificare valori di significatività inferiori o uguali al livello più grande.
- Se si seleziona **Utilizza sottoscritti in stile APA**, il secondo valore viene ignorato.

Adatta valori P per confronti multipli

La correzione **Bonferroni** regola per il tasso di errore di famiglia (FWER). Il metodo **Benjamini-Hochberg** è un errore di regolazione FDR (false discovery rate). Questo metodo è meno conservatore rispetto alla correzione Bonferroni.

Test di indipendenza (Chi-quadrato)

Chi - quadrato prova di indipendenza per i tavoli in cui esiste almeno una variabile di categoria sia nelle righe che nelle colonne.

Usa i totali parziali anziché le categorie subtotalizzate

Ogni subtotale sostituisce le sue categorie per i test di significatività. In caso contrario, solo i subtotali per i quali le categorie subtotali sono nascoste sostituiscono le loro categorie per i test.

Includi variabili a risposta multipla nei test

Le categorie di serie di risposta multipla sono incluse nei test di significatività. In caso contrario, i set di risposta multipla non sono inclusi nei test di significatività.

Tabelle semplici per Variabili categoriali

Tabelle semplici per Variabili categoriali

La maggior parte delle tabelle che si desidera creare includerà probabilmente almeno una **variabile categoriale**. Una variabile categoriale è una con un numero limitato di valori distinti o categorie (per esempio sesso o religione). Le variabili categoriali possono essere **nominali** o **ordinali**.

- **Nominale**. Una variabile può essere considerata nominale quando i suoi valori rappresentano categorie prive di classificazione intrinseca (ad esempio, il dipartimento della società in cui lavora un dipendente). Degli esempi di variabili nominali includono la regione, il codice postale e l'affiliazione religiosa.
- **Ordinale**. Una variabile può essere trattata come ordinale quando i suoi valori rappresentano categorie con una classificazione intrinseca (ad esempio, livelli di soddisfazione del servizio da molto insoddisfatti a molto soddisfatti). Degli esempi di variabili ordinali includono i punteggi di atteggiamento che rappresentano i gradi di soddisfazione o fiducia e i punteggi di classificazione delle preferenze.

Un'icona accanto ad ogni variabile nell'elenco delle variabili identifica il tipo di variabile.

Tabella 7. Icone del livello di misurazione				
	Numerico	Stringa	Dati	Ora
Scala (continuo)		n/d		
Ordinale				
Nominale				

Le tabelle personalizzate sono ottimizzate per l'utilizzo con variabili categoriali che hanno definito **etichette di valore**. Per ulteriori informazioni, consultare l'argomento [“Tabelle di costruzione”](#) a pagina 1.

File di dati di esempio

Gli esempi in questo capitolo utilizzare il file di `datisurvey_sample.sav`. Per ulteriori informazioni, consultare l'argomento [“File di esempio”](#) a pagina 76.

Tutti gli esempi forniti qui visualizzano etichette variabili in finestre di dialogo, ordinate in ordine alfabetico. Le proprietà di visualizzazione della lista variabili sono impostate sulla scheda Generale nella finestra di dialogo Opzioni (Modifica menu, Opzioni).

Una variabile singola categoriale

Anche se una tabella di una singola variabile categoriale può essere una delle tabelle più semplici che si possono creare, potrebbe essere spesso tutto ciò che si desidera o di cui si ha bisogno.

1. Dai menu, scegliere:

Analisi > Tavole > Tabelle personalizzate ...

2. Nel builder table, trascinare e rilasciare la *Categoria Età* dall'elenco delle variabili all'area Rows sul riquadro canvas.

Un'anteprima della tabella viene visualizzata sul riquadro canvas. L'anteprima non visualizza valori dati reali; visualizza solo i segnaposto dove verranno visualizzati i dati.

3. Fare clic su **OK** per creare la tabella.

La tabella viene visualizzata nella finestra del Viewer.

		Count
Age category	Less than 25	242
	25 to 34	627
	35 to 44	679
	45 to 54	481
	55 to 64	320
	65 or older	479

Figura 2. Variabile categoriale singola in righe

In questa semplice tabella, l'intestazione della colonna *Conteggio* non è davvero necessaria e si può creare la tabella senza questa intestazione di colonna.

4. Aprire di nuovo il builder table (Analizza menu, tabelle, Tabelle personalizzate).
5. Nel gruppo di statistiche di riepilogo, selezionare (clicca) **Nascondi** per Posizione.
6. Fare clic su **OK** per creare la tabella.

Age category	Less than 25	242
	25 to 34	627
	35 to 44	679
	45 to 54	481
	55 to 64	320
	65 or older	479

Figura 3. Variabile categoriale singola senza etichetta di colonna statistiche di riepilogo

Percentuali

Oltre ai conteggi, è possibile anche visualizzare le percentuali. Per una semplice tabella di una singola variabile categoriale, se la variabile viene visualizzata in righe, probabilmente si desidera guardare le percentuali di colonna. Viceversa, per una variabile visualizzata in colonne, probabilmente si desidera guardare le percentuali di riga.

1. Aprire di nuovo il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Nel gruppo di statistiche di riepilogo, deselezionare **Nascondi** per Posizione. Dal momento che questa tabella avrà due colonne, si desidera visualizzare le etichette delle colonne in modo da sapere cosa rappresenta ogni colonna.
3. Fare clic con il tasto destro del mouse su *Age category* sul riquadro canvas e selezionare **Statistiche di riepilogo** dal menu a comparsa.
4. Nella finestra di dialogo Statistiche di riepilogo, selezionare **Colonna N%** nell'elenco Statistiche e fare clic sulla freccia per aggiungendola alla lista Visualizza.
5. Nella cella Etichetta nell'elenco Visualizza, eliminare l'etichetta predefinita e immettere Percent.
6. Fare clic su **Applica alla selezione** e quindi fare clic su **OK** nel builder da tavolo per creare la tabella.

		Count	Percent
Age category	Less than 25	242	8.6%
	25 to 34	627	22.2%
	35 to 44	679	24.0%
	45 to 54	481	17.0%
	55 to 64	320	11.3%
	65 or older	479	16.9%

Figura 4. Percentuali di conteggi e colonne

Totali

I totali non vengono inclusi automaticamente nelle tabelle personalizzate, ma è facile aggiungere totali a una tabella.

1. Aprire di nuovo il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Fare clic con il tasto destro del mouse su *Age category* sul riquadro canvas e selezionare **Categorie e Totali** dal menu a comparsa.
3. Selezionare (cliccare) **Totale** nella finestra di dialogo Categorie e Totali.
4. Fare clic su **Applica** e quindi fare clic su **OK** nel builder da tavolo per creare la tabella.

		Count	Percent
Age category	Less than 25	242	8.6%
	25 to 34	627	22.2%
	35 to 44	679	24.0%
	45 to 54	481	17.0%
	55 to 64	320	11.3%
	65 or older	479	16.9%
Total		2828	100.0%

Figura 5. Conteggi, percentuali di colonna e totali

Per ulteriori informazioni, consultare l'argomento [“Totali e totali per Variabili categoriali”](#) a pagina 30.

Tavola di contingenza

Crosstabula è una tecnica di base per esaminare il rapporto tra due variabili categoriali. Ad esempio, utilizzando la *Categoria Età* come variabile di riga e *Gender* come variabile di colonna, è possibile creare una crosstabulazione bidimensionale che mostra il numero di maschi e femmine in ogni categoria di età.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Clicca su **Reset** per eliminare eventuali selezioni precedenti nel builder table.
3. Nel builder table, trascinare e rilasciare la *Categoria Età* dall'elenco delle variabili all'area Rows sul riquadro canvas.
4. Trascinare e rilasciare *Gender* dall'elenco delle variabili all'area Colonne sul riquadro canvas. (Potrebbe essere necessario scorrere verso il basso attraverso l'elenco delle variabili per trovare questa variabile.)
5. Fare clic su **OK** per creare la tabella.

		Gender	
		Male	Female
		Count	Count
Age category	Less than 25	108	134
	25 to 34	276	351
	35 to 44	309	370
	45 to 54	221	260
	55 to 64	136	184
	65 or older	178	301

Figura 6. Crosstabulazione della categoria Age e Gender

Percentuali in Crosstabulazioni

In una crosta bidimensionale, le percentuali di fila e di colonna possono fornire informazioni utili.

1. Aprire di nuovo il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Fare clic con il tasto destro del mouse su *Gender* sul riquadro canvas.

Si può notare che **Statistiche di riepilogo** è disabilitata nel menu a comparsa. Questo perché è possibile selezionare statistiche di riepilogo solo per la variabile innerpiù nella dimensione di origine statistiche. La dimensione di origine delle statistiche predefinite (riga o colonna) per le variabili categoriali si basa sull'ordine in cui si trascina e si abbassano le variabili sul riquadro canvas. In questo esempio, abbiamo trascinato la *Categoria Età* alla dimensione delle righe prima -- e dato che non ci sono altre variabili nella dimensione delle righe, *Categoria Età* è la variabile di origine statistica. È possibile modificare la dimensione di origine delle statistiche, ma in questo esempio non è necessario farlo. Per ulteriori informazioni, consultare l'argomento [“Statistiche di riepilogo”](#) a pagina 6.

3. Fare clic con il tasto destro del mouse su *Age category* sul riquadro canvas e selezionare **Statistiche di riepilogo** dal menu a comparsa.
4. Nella finestra di dialogo Statistiche di riepilogo, selezionare **Colonna N%** nell'elenco Statistiche e fare clic sulla freccia per aggiungendola alla lista Visualizza.
5. Selezionare **Riga N%** nell'elenco Statistiche e fare clic sulla freccia per aggiungendola alla lista Visualizza.
6. Fare clic su **Applica alla selezione** e quindi fare clic su **OK** nel builder da tavolo per creare la tabella.

		Gender					
		Male			Female		
		Count	Column N %	Row N %	Count	Column N %	Row N %
Age category	Less than 25	108	8.8%	44.6%	134	8.4%	55.4%
	25 to 34	276	22.5%	44.0%	351	21.9%	56.0%
	35 to 44	309	25.2%	45.5%	370	23.1%	54.5%
	45 to 54	221	18.0%	45.9%	260	16.3%	54.1%
	55 to 64	136	11.1%	42.5%	184	11.5%	57.5%
	65 or older	178	14.5%	37.2%	301	18.8%	62.8%

Figura 7. Crosstabulazione con percentuali di fila e colonne

Controllo formato di visualizzazione

È possibile controllare il formato di visualizzazione, incluso il numero di decimali visualizzati nelle statistiche di riepilogo. Ad esempio, per impostazione predefinita, le percentuali vengono visualizzate con un decimale e un segno per cento. Ma se volete che i valori cellulari mostrino due decimali e nessun segno per cento?

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Fare clic con il tasto destro del mouse su *Age category* sul riquadro canvas e selezionare **Statistiche di riepilogo** dal menu a comparsa.
3. Per le due statistiche di riepilogo percentuale selezionate (**% N colonna** e **% N riga**), selezionare **nnnn.n** dall'elenco a discesa Formato e immettere 2 nella cella Decimali per entrambi.
4. Fare clic su **OK** per creare la tabella.

		Gender					
		Male			Female		
		Count	Column N %	Row N %	Count	Column N %	Row N %
Age category	Less than 25	108	8.79	44.63	134	8.38	55.37
	25 to 34	276	22.48	44.02	351	21.94	55.98
	35 to 44	309	25.16	45.51	370	23.13	54.49
	45 to 54	221	18.00	45.95	260	16.25	54.05
	55 to 64	136	11.07	42.50	184	11.50	57.50
	65 or older	178	14.50	37.16	301	18.81	62.84

Figura 8. Visualizzazione delle celle formattate per le percentuali di riga e colonna

Totali marginali

E'abbastanza comune nelle crosstabulazioni per visualizzare **totali marginali**-- totali per ogni riga e colonna. Poiché questi non sono inclusi nelle tabelle personalizzate per impostazione predefinita, è necessario aggiungerli esplicitamente alle tabelle.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Clicca su **Reset** per eliminare eventuali selezioni precedenti nel builder table.
3. Nel builder table, trascinare e rilasciare la *Categoria Età* dall'elenco delle variabili all'area Rows sul riquadro canvas.
4. Trascinare e rilasciare *Gender* dall'elenco delle variabili all'area Colonne sul riquadro canvas. (Potrebbe essere necessario scorrere verso il basso attraverso l'elenco delle variabili per trovare questa variabile.)
5. Fare clic con il tasto destro del mouse su *Age category* sul riquadro canvas e selezionare **Categorie e Totali** dal menu a comparsa.
6. Selezionare (cliccare) **Totale** nella finestra di dialogo Categorie e Totali e quindi fare clic su **Applica**.
7. Fare clic con il tasto destro del mouse su *Gender* sul riquadro canvas e selezionare **Categorie e Totali** dal menu a comparsa.
8. Selezionare (cliccare) **Totale** nella finestra di dialogo Categorie e Totali e quindi fare clic su **Applica**.
9. Nel gruppo di statistiche di riepilogo, selezionare (clicca) **Nascondi** per Posizione. (Visto che stai visualizzando solo conteggi, non è necessario identificare la "statistica" visualizzata nelle celle dati della tabella.)
10. Fare clic su **OK** per creare la tabella.

		Gender		
		Male	Female	Total
Age category	Less than 25	108	134	242
	25 to 34	276	351	627
	35 to 44	309	370	679
	45 to 54	221	260	481
	55 to 64	136	184	320
	65 or older	178	301	479
	Total	1228	1600	2828

Figura 9. Crosstabulazione con totali marginali

Ordinamento e categorie escluse

Per impostazione predefinita, le categorie vengono visualizzate nell'ordine crescente dei valori dei dati che le etichette dei valori di categoria rappresentano. Ad esempio, sebbene le etichette di valore di *Meno di 25*, *da 25 a 34*, *da 35 a 44*, ..., ecc., sono visualizzati per le categorie di età, i valori effettivi di dati sottostanti sono 1, 2, 3, ..., ecc., ed è proprio quei valori di dati sottostanti che controllano l'ordine di visualizzazione predefinito delle categorie.

È possibile modificare facilmente l'ordine delle categorie e anche escludere le categorie che non si desidera vengano visualizzate nella tabella.

ordinamento di categorie

È possibile riorganizzare manualmente le categorie o le categorie di ordinamento in ordine crescente o decrescente di:

- Valori dei dati.
 - Etichette di valore.
 - Conteggi celle.
 - Statistiche di riepilogo. Le statistiche di riepilogo disponibili per l'ordinamento dipendono dalle statistiche di riepilogo selezionate per visualizzare in tabella.
1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
 2. Se *Age category* non è già visualizzato nell'area Rows sul riquadro canvas, trascinare e rilasciarlo lì.
 3. Fare clic con il tasto destro del mouse su *Age category* sul riquadro canvas e selezionare **Categorie e Totali** dal menu a comparsa.
- Entrambi i valori dei dati e le etichette di valore associate sono visualizzati nell'ordine di ordinamento corrente, che in questo caso è ancora ascendente di valori dei dati.
4. Nel gruppo Categorie di ordinamento, selezionare **Descending** dall'elenco a discesa dell'ordine.
- Il criterio di ordinamento viene invertito.
5. Selezionare **Labels** dall'elenco a discesa By.

Le categorie sono ora ordinate in ordine alfabetico discendente delle etichette di valore.

Si noti che la categoria etichettata *Meno di 25* è in cima alla lista. In ordine alfabetico le lettere arrivano dopo i numeri. Dal momento che questa è l'unica etichetta che inizia con una lettera e dato che l'elenco è ordinato in ordine decrescente (reverse), questo tipo di categoria in cima alla lista.

Se si desidera che una determinata categoria appaia in una posizione diversa nella lista, è possibile spostarla facilmente.

6. Fare clic sulla categoria etichettata *Meno di 25* nella lista Label.
7. Clicca sulla freccia in basso a destra dell'elenco. La categoria si sposta in basso di una riga nell'elenco.
8. Continua a clicca sulla freccia verso il basso fino a quando la categoria non è in fondo alla lista.

esclusione di categorie

Se ci sono alcune categorie che non si desidera comparire nella tabella, è possibile escluderle.

1. Fare clic sulla categoria etichettata *Meno di 25* nella lista Label.
2. Fare clic sul tasto freccia a sinistra dell'elenco Escludi.
3. Fare clic sulla categoria etichettata *65 o older* nella lista Label.
4. Clicca nuovamente il tasto freccia a sinistra della lista Escludi.

Le due categorie vengono spostate dall'elenco Visualizza alla lista Escludi. Se si cambia idea, è possibile spostarle facilmente nella lista Visualizza.

5. Fare clic su **Applica** e quindi fare clic su **OK** nel builder da tavolo per creare la tabella.

		Gender		
		Male	Female	Total
Age category	55 to 64	136	184	320
	45 to 54	221	260	481
	35 to 44	309	370	679
	25 to 34	276	351	627
	Total	942	1165	2107

Figura 10. Tabella ordinata per etichetta del valore decrescente, alcune categorie escluse

Si noti che i totali sono inferiori a quelli che erano prima che le due categorie fossero escluse. Questo perché i totali si basano sulle categorie incluse nella tabella. Eventuali categorie escluse sono escluse dal calcolo totale. Per ulteriori informazioni, consultare l'argomento [“Totali e totali per Variabili categoriali”](#) a pagina 30.

Stack, Nesting e Layers con Variabili categoriali

Stack, nidificazione e strati sono tutti metodi per la visualizzazione di più variabili nella stessa tabella. Questo capitolo si concentra sull'utilizzo di queste tecniche con variabili categoriali, sebbene possano essere utilizzate anche con variabili di scala.

File di dati di esempio

Gli esempi in questo capitolo utilizzano il file di dati *survey_sample.sav*. Per ulteriori informazioni, consultare l'argomento [“File di esempio”](#) a pagina 76.

Tutti gli esempi forniti qui visualizzano etichette variabili in finestre di dialogo, ordinate in ordine alfabetico. Le proprietà di visualizzazione della lista variabili sono impostate sulla scheda Generale nella finestra di dialogo Opzioni (Modifica menu, Opzioni).

Variabili Categoriali categoriali

Stacking può essere pensato come prendere tavole separate e incollarle insieme nello stesso display. Ad esempio, è possibile visualizzare le informazioni su *Gender* e *Categoria Età* in sezioni separate della stessa tabella.

1. Dai menu, scegliere:

Analisi > Tavole > Tabelle personalizzate ...

2. Nel builder table, trascinare e rilasciare *Gender* dall'elenco variabile all'area Rows sul riquadro canvas.
3. Trascinare e rilasciare la *Categoria Età* dall'elenco delle variabili all'area Rows sotto *Gender*.

Le due variabili sono ora impilate nella dimensione della riga.

4. Fare clic su **OK** per creare la tabella.

		Count
Gender	Male	1232
	Female	1600
Age category	Less than 25	242
	25 to 34	627
	35 to 44	679
	45 to 54	481
	55 to 64	320
	65 or older	479

Figura 11. Tabella delle variabili categoriali impilate in righe

È possibile anche pilare variabili in colonne in una moda simile.

Impilamento con Crosstabulazione

Una tabella in sovrapposizione può includere altre variabili in altre dimensioni. Ad esempio, si potrebbero crogiare due variabili impilate nelle righe con una terza variabile visualizzata nella dimensione della colonna.

1. Aprire il builder table di nuovo (Analizza menu, tabelle, Tabelle personalizzate).
2. Se *Age category* e *Gender* non sono già impilati nelle righe, seguire le indicazioni sopra per impilarle.
3. Trascinare e rilasciare *Otteni notizie da internet* dall'elenco delle variabili all'area Colonne sul riquadro canvas.
4. Fare clic su **OK** per creare la tabella.

		Get news from internet	
		No	Yes
		Count	Count
Gender	Male	873	359
	Female	1092	508
Age category	Less than 25	146	96
	25 to 34	368	259
	35 to 44	435	244
	45 to 54	346	135
	55 to 64	252	68
	65 or older	416	63

Figura 12. Due variabili di riga impilate crosstabulate con una variabile di colonna

Nota: ci sono diverse variabili con etichette che iniziano con *Get news da ...*, quindi può essere difficile distinguerli nell'elenco delle variabili (dato che le etichette possono essere troppo largate per essere visualizzate completamente nell'elenco delle variabili). Ci sono due modi per vedere l'intera etichetta variabile:

- Posizionare il puntatore del mouse su una variabile nell'elenco per visualizzare l'intera etichetta in un ToolTipa comparsa.
- Clicca e trascina la barra verticale che separa le liste variabili e Categorie dal riquadro canvas per rendere le liste più ampie.

Variabili Categoriali Di Nidificazione

Nesting, come la crosstabulazione, può mostrare il rapporto tra due variabili categoriali, tranne che una variabile è nidificata all'interno dell'altra nella stessa dimensione. Ad esempio, si potrebbe nidificare *Gender* all'interno di *Categoria di età* nella dimensione della riga, mostrando il numero di maschi e femmine in ogni categoria di età.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Clicca su **Reset** per eliminare eventuali selezioni precedenti nel builder table.
3. Nel builder table, trascinare e rilasciare la *Categoria Età* dall'elenco delle variabili all'area Rows sul riquadro canvas.
4. Trascinare e rilasciare *Gender* dall'elenco delle variabili a destra della *Categoria Età* nell'area Rows.

L'anteprima sul riquadro canvas mostra ora che la tabella nidificata conterrà una singola colonna di conteggi, con ogni cella contenente il numero di maschi o femmine in ogni categoria di età.

Si può notare che l'etichetta variabile *Gender* viene visualizzata ripetutamente, una volta per ogni categoria di età. È possibile ridurre al minimo questo tipo di ripetizione posizionando la variabile con le categorie fewest a livello più esterno della nidificazione.

5. Clicca sull'etichetta variabile *Gender* sul riquadro canvas.

6. Trascinare e lasciar cadere la variabile fino a sinistra nella zona di Rassi come si può.

Ora invece di *Gender* ripetuto sei volte, *Categoria Età* si ripete due volte. Si tratta di un tavolo meno clitante che produrrà essenzialmente gli stessi risultati.

7. Fare clic su **OK** per creare la tabella.

				Count
Gender	Male	Age category	Less than 25	108
			25 to 34	276
			35 to 44	309
			45 to 54	221
			55 to 64	136
			65 or older	178
	Female	Age category	Less than 25	134
			25 to 34	351
			35 to 44	370
			45 to 54	260
			55 to 64	184
			65 or older	301

Figura 13. Categoria della categoria Età nidificata all'interno di Gender

Nota: Tavole personalizzate non onorificano l'elaborazione del file suddiviso a strati. Per ottenere lo stesso risultato dei file split stratificati, posizionare le variabili del file suddiviso negli strati di nidificazione più esterni della tabella.

Soppressione Etichette variabili

Un'altra soluzione alle etichette variabili ridondanti in tabelle nidificate è semplicemente quella di sopprimere la visualizzazione di nomi o etichette variabili. Dal momento che le etichette di valore sia per *Gender* che per *Age category* sono probabilmente sufficientemente descrittive senza le etichette variabili, possiamo eliminare le etichette per entrambe le variabili.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).

2. Fare clic con il tasto destro del mouse su *Age category* sul riquadro canvas e deselezionare **Mostra Label variabile** sul menu a comparsa.

3. Fare lo stesso per *Gender*.

Le etichette variabili sono ancora visualizzate nell'anteprima della tabella, ma non saranno incluse nella tabella.

4. Fare clic su **OK** per creare la tabella.

		Count
Male	Less than 25	108
	25 to 34	276
	35 to 44	309
	45 to 54	221
	55 to 64	136
	65 or older	178
Female	Less than 25	134
	25 to 34	351
	35 to 44	370
	45 to 54	260
	55 to 64	184
	65 or older	301

Figura 14. Tabella nidificata senza etichette variabili

Se volete le etichette variabili incluse con la tabella da qualche parte -- senza visualizzarle più volte nel corpo della tabella - potete includerle nel titolo della tabella o nell'etichetta d'angolo.

5. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
6. Fare clic sulla scheda **Titoli**.
7. Clicca ovunque nella casella di testo Titolo.
8. Clicca su **Expression di tabella**. Il testo & [*Expression Table*] viene visualizzato nella casella di testo Titolo. In questo modo verrà generato un titolo di tabella che include le etichette variabili per le variabili utilizzate nella tabella.
9. Fare clic su **OK** per creare la tabella.

Gender > Age category

		Count
Male	Less than 25	108
	25 to 34	276
	35 to 44	309
	45 to 54	221
	55 to 64	136
	65 or older	178
Female	Less than 25	134
	25 to 34	351
	35 to 44	370
	45 to 54	260
	55 to 64	184
	65 or older	301

Figura 15. Etichette variabili in titolo da tavolo

Il maggiore del segno (>) nel titolo indica che *Categoria Età* è nidificata all'interno di *Gender*.

Tavole di contingenza nidificate

Una tabella nidificata può contenere altre variabili in altre dimensioni. Ad esempio, è possibile nidificare *Categoria Età* all'interno di *Gender* nelle righe e crosticamente le righe nidificate con una terza variabile nella dimensione della colonna.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Se *Age category* non è già nidificato all'interno di *Gender* nelle righe, seguire le indicazioni riportate nell'esempio precedente per nestarle.
3. Trascinare e rilasciare *Ottieni notizie da internet* dall'elenco delle variabili all'area Colonne sul riquadro canvas.

Si può notare che la tabella è troppo grande per visualizzare completamente sul riquadro canvas. È possibile scorrere su / giù o destra / sinistra sul riquadro canvas per vedere più in anteprima della tabella, oppure è possibile:

- Clicca su **Compact** nel builder da tavolo per visualizzare una vista compatta. In questo modo vengono visualizzate solo le etichette variabili, senza alcuna informazione sulle categorie o le statistiche di riepilogo incluse nella tabella.
- Aumentare la dimensione del builder da tavolo cliccando e trascinando qualsiasi lato o angoli del costruttore di tabelle.

4. Fare clic su **OK** per creare la tabella.

				Get news from internet	
				No	Yes
				Count	Count
Gender	Male	Age category	Less than 25	59	49
			25 to 34	159	117
			35 to 44	217	92
			45 to 54	169	52
			55 to 64	112	24
			65 or older	155	23
	Female	Age category	Less than 25	87	47
			25 to 34	209	142
			35 to 44	218	152
			45 to 54	177	83
			55 to 64	140	44
			65 or older	261	40

Figura 16. Tavole di contingenza nidificate

Scambio di righe e colonne

Cosa si fa se si passa molto tempo a configurare un tavolo complesso e poi a decidere è assolutamente perfetto - tranne che si vuole commutare l'orientamento, mettendo tutte le variabili di riga nelle colonne e viceversa? Ad esempio, hai creato una crosstabulazione nidificata con *Categoria Età* e *Gender* nidificata nelle righe, ma ora si desidera che queste due variabili demografiche nidificino nelle colonne.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Fare clic con il tasto destro del mouse ovunque sul riquadro canvas e selezionare **Swap Row e Variabili di colonna** dal menu a comparsa.

Le variabili di riga e colonna sono state ora commutate.

Prima di creare la tabella, facciamo qualche modifica per rendere il display meno cliccettato.

3. Selezionare **Nascondi** per sopprimere la visualizzazione dell'etichetta di colonna delle statistiche di riepilogo.
4. Fare clic con il tasto destro del mouse su *Gender* sul riquadro canvas e deselezionare **Mostra Label variabile**.
5. Ora clicca su **OK** per creare la tabella.

		Male						Female					
		Age category						Age category					
		Less than 25	25 to 34	35 to 44	45 to 54	55 to 64	65 or older	Less than 25	25 to 34	35 to 44	45 to 54	55 to 64	65 or older
Get news from internet	No	59	159	217	169	112	155	87	209	218	177	140	261
	Yes	49	117	92	52	24	23	47	142	152	83	44	40

Figura 17. Crosstabulazione con variabili demografiche nidificate nelle colonne

Livelli

È possibile utilizzare gli strati per aggiungere una dimensione di profondità alle proprie tabelle, creando cubi tridimensionali ". Gli strati sono, infatti, abbastanza simili a nidificare o impilare; la differenza primaria è che solo una categoria di strati è visibile alla volta. Ad esempio, l'utilizzo di *Categoria Età* come variabile di riga e *Gender* come variabile di layer produce una tabella in cui vengono visualizzate informazioni per maschi e femmine in diversi strati della tabella.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Clicca su **Reset** per eliminare eventuali selezioni precedenti nel builder table.
3. Nel builder table, trascinare e rilasciare la *Categoria Età* dall'elenco delle variabili all'area Rows sul riquadro canvas.
4. Clicca su **Layers** nella parte superiore del builder da tavolo per visualizzare l'elenco Layers.
5. Trascinare e rilasciare *Gender* dall'elenco delle variabili alla lista dei Layers.

A questo punto si potrebbe notare che l'aggiunta di una variabile di layer non ha alcun effetto visibile sull'anteprima visualizzata sul riquadro canvas. Le variabili di strato non influenzano l'anteprima sul riquadro canvas a meno che la variabile di layer non sia la variabile di origine statistica e si modificano le statistiche di riepilogo.

6. Fare clic su **OK** per creare la tabella.

Gender Male		Count
Age category	Less than 25	108
	25 to 34	276
	35 to 44	309
	45 to 54	221
	55 to 64	136
	65 or older	178

Figura 18. Tabella semplice stratificata

A prima vista, questo tavolo non sembra diverso da un semplice tavolo di una singola variabile categoriale. L'unica differenza è la presenza dell'etichetta *Genere Maschile* in cima alla tabella.

7. Fare doppio clic sulla tabella nella finestra del Viewer per attivarla.
8. È ora possibile vedere che l'etichetta *Genere Maschile* è in realtà una scelta in un elenco a discesa.
9. Clicca sulla freccia in basso sull'elenco a discesa per visualizzare l'intero elenco degli strati.

In questa tabella c'è solo un'altra scelta nella lista.

10. Selezionare *Gender Female* dall'elenco a discesa.

Gender Female		Count
Age category	Less than 25	134
	25 to 34	351
	35 to 44	370
	45 to 54	260
	55 to 64	184
	65 or older	301

Figura 19. Semplice tavola stratificata con strato diverso visualizzato

Due Variabili di strato categoriale sovrapposte

Se si dispone di più di una variabile categoriale negli strati, è possibile pilare o nidificare le variabili di strato. Per impostazione predefinita, le variabili di layer sono impilate. (Nota: se si hanno delle variabili di layer di scala, le variabili di strato possono essere solo impilate.)

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Se non si dispone già di *Categoria Età* nelle righe e *Gender* negli strati, seguire le indicazioni nell'esempio precedente per la creazione di una tabella letta.
3. Trascinare e rilasciare *Altezza più alta* dall'elenco delle variabili all'elenco Layer sotto *Gender*.

I due pulsanti radio sotto l'elenco Layer nel gruppo Layer Output sono ora attivati. La selezione predefinita è **Mostra ogni categoria come uno strato**. Questo equivale ad impilare.

4. Fare clic su **OK** per creare la tabella.
5. Fare doppio clic sulla tabella nella finestra del Viewer per attivarla.
6. Clicca sulla freccia in basso sull'elenco a discesa per visualizzare l'intero elenco degli strati.

Nella tabella sono presenti sette strati: due strati per le due categorie *Gender* e cinque strati per le cinque categorie *Highest degree*. Per gli strati sovrapposti, il numero totale di layer è la somma del numero di categorie per le variabili di strato (incluse le eventuali categorie totali o sottototali richieste per le variabili di strato).

Due Variabili di strato categoriale nidificato

Le variabili di strato categoriale di nidificazione crea uno strato separato per ogni combinazione di categorie di variabili di strato.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Se non l'avete già fatto, seguite le indicazioni nell'esempio precedente per la creazione di una tabella di strati impilati.
3. Nel gruppo di Output Layer selezionare **Mostra ogni combinazione di categorie come uno strato**. Questo equivale a nidificare.
4. Fare clic su **OK** per creare la tabella.
5. Fare doppio clic sulla tabella nella finestra del Viewer per attivarla.
6. Clicca sulla freccia in basso sull'elenco a discesa per visualizzare l'intero elenco degli strati.

Ci sono 10 strati nella tabella (bisogna scorrere la lista per vederli tutti), uno per ogni combinazione di *Gender* e *Highest degree*. Per gli strati nidificati, il numero totale di strati è il *prodotto* del numero di categorie per ogni variabile di strato (in questo esempio, $5 \times 2 = 10$).

Stampa delle tabelle Layered

Per impostazione predefinita, viene stampato solo lo strato attualmente visibile. Per stampare tutti gli strati di una tabella:

1. Fare doppio clic sulla tabella nella finestra del Viewer per attivarla.
2. Dalla finestra Viewer , scegliere:

Formato > Proprietà tabella...

3. Fare clic sulla scheda **Stampa**.
4. Selezionare **Stampa tutti gli strati**.

È inoltre possibile salvare questa impostazione come parte di un sito TableLook,, compreso il sito predefinito TableLook.

Totali e totali per Variabili categoriali

È possibile includere entrambi i totali e i subtotali nelle tabelle personalizzate. I totali e i subtotali possono essere applicati a variabili categoriali a qualsiasi livello di nidificazione in qualsiasi dimensione - riga, colonna o strato.

File di dati di esempio

Gli esempi in questo capitolo utilizzano il file di dati *survey_sample.sav*. Per ulteriori informazioni, consultare l'argomento [“File di esempio”](#) a pagina 76.

Tutti gli esempi forniti qui visualizzano etichette variabili in finestre di dialogo, ordinate in ordine alfabetico. Le proprietà di visualizzazione della lista variabili sono impostate sulla scheda Generale nella finestra di dialogo Opzioni (Modifica menu, Opzioni).

Totale semplice per una singola variabile

1. Dai menu, scegliere:
Analisi > Tavole > Tabelle personalizzate ...
2. Nel builder table, trascinare e rilasciare la *Categoria Età* dall'elenco delle variabili all'area Rows sul riquadro canvas.
3. Fare clic con il tasto destro del mouse su *Age category* sul riquadro canvas e scegliere **Statistiche di riepilogo** dal menu a comparsa.
4. Nella finestra di dialogo Statistiche di riepilogo, selezionare **Colonna N%** nell'elenco Statistiche e fare clic sulla freccia per aggiungendola alla lista Visualizza.

5. Nella cella Etichetta nell'elenco Visualizza, eliminare l'etichetta predefinita e immettere Percent.
6. Fare clic su **Applica alla selezione**.
7. Fare clic con il tasto destro del mouse su *Age category* sul riquadro canvas e scegliere **Categorie e Totali** dal menu a comparsa.
8. Selezionare (cliccare) **Totale** nella finestra di dialogo Categorie e Totali.
9. Fare clic su **Applica** e quindi fare clic su **OK** nel builder da tavolo per creare la tabella.

		Count	Percent
Age category	Less than 25	242	8.6%
	25 to 34	627	22.2%
	35 to 44	679	24.0%
	45 to 54	481	17.0%
	55 to 64	320	11.3%
	65 or older	479	16.9%
	Total	2828	100.0%

Figura 20. Totale semplice per una singola variabile categoriale

Ciò Che Vedete È Quello Che Viene Totalizzato

I totali si basano sulle categorie visualizzate nella tabella. Se si sceglie di escludere alcune categorie da una tabella, i casi da quelle categorie non sono inclusi nei calcoli totali.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Fare clic con il tasto destro del mouse su *Age category* sul riquadro canvas e scegliere **Categorie e Totali** dal menu a comparsa.
3. Fare clic sulla categoria etichettata *Meno di 25* nella lista Label.
4. Fare clic sul tasto freccia a sinistra dell'elenco Escludi.
5. Fare clic sulla categoria etichettata *65 o older* nella lista Label.
6. Clicca nuovamente il tasto freccia a sinistra della lista Escludi.

Le due categorie vengono spostate dall'elenco Visualizza alla lista Escludi.

7. Fare clic su **Applica** e quindi fare clic su **OK** nel builder da tavolo per creare la tabella.

		Count	Percent
Age category	25 to 34	627	29.8%
	35 to 44	679	32.2%
	45 to 54	481	22.8%
	55 to 64	320	15.2%
	Total	2107	100.0%

Figura 21. Totale in tavola con categorie escluse

Il conteggio totale in questa tabella è di soli 2.107, rispetto al 2.828 quando tutte le categorie sono incluse. Solo le categorie che vengono utilizzate nella tabella sono incluse nel totale. (Il totale percentuale è ancora del 100% perché tutte le percentuali si basano sul totale dei casi utilizzati nella tabella, non sul totale dei casi nel file dei dati.)

Visualizza posizione di Totali

Per impostazione predefinita, i totali vengono visualizzati al di sotto delle categorie che si stanno totalizzando. È possibile modificare la posizione di visualizzazione dei totali per mostrarli al di sopra delle categorie che si stanno totalizzando.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Fare clic con il tasto destro del mouse su *Age category* sul riquadro canvas e scegliere **Categorie e Totali** dal menu a comparsa.
3. Nel gruppo Totali e Sottogruppi Aspetto selezionare **Sopra categorie a cui si applicano**.

4. Fare clic su **Applica** e quindi fare clic su **OK** nel builder da tavolo per creare la tabella.

		Count	Percent
Age category	Total	2107	100.0%
	25 to 34	627	29.8%
	35 to 44	679	32.2%
	45 to 54	481	22.8%
	55 to 64	320	15.2%

Figura 22. Totale visualizzato sopra le categorie totali

Totale per le tabelle Nestate

Poiché i totali possono essere applicati a variabili categoriali a qualsiasi livello di nidificazione, è possibile creare tabelle che contengono totali di gruppo a più livelli di nidificazione.

Totale di gruppo

I totali per le variabili categoriali nidificate all'interno di altre variabili categoriali rappresentano i totali di gruppo.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Trascinare e rilasciare *Gender* a sinistra della *Categoria Età* sul riquadro canvas.
3. Fare clic con il tasto destro del mouse su *Age category* sul riquadro canvas e scegliere **Categorie e Totali** dal menu a comparsa.

Prima di creare la tabella, spostiamo i totali al di sotto delle categorie totali.

4. Nel gruppo Totali e Sottogruppi Aspetto, selezionare **Categorie di Below a cui si applicano**.
5. Clicca su **Applica** per salvare l'impostazione e tornare al builder da tavolo.
6. Fare clic su **OK** per creare la tabella.

				Count	Percent
Gender	Male	Age category	25 to 34	276	29.3%
			35 to 44	309	32.8%
			45 to 54	221	23.5%
			55 to 64	136	14.4%
			Total	942	100.0%
	Female	Age category	25 to 34	351	30.1%
			35 to 44	370	31.8%
			45 to 54	260	22.3%
			55 to 64	184	15.8%
			Total	1165	100.0%

Figura 23. Fasce di età totali all'interno delle categorie Gender

La tabella visualizza ora due totali di gruppo: uno per i maschi e uno per le femmine.

totali finali

I totali applicati alle variabili nidificate sono sempre totali di gruppo, non totali. Se si desidera totali per l'intera tabella, è possibile applicare totali alla variabile a livello di nidificazione più esterno.

1. Aprire di nuovo il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Fare clic con il tasto destro del mouse su *Gender* sul riquadro canvas e scegliere **Categorie e Totali** dal menu a comparsa.
3. Selezionare (cliccare) **Totale** nella finestra di dialogo Categorie e Totali.
4. Fare clic su **Applica** e quindi fare clic su **OK** nel builder da tavolo per creare la tabella.

				Count	Percent
Gender	Male	Age category	25 to 34	276	29.3%
			35 to 44	309	32.8%
			45 to 54	221	23.5%
			55 to 64	136	14.4%
			Total	942	100.0%
	Female	Age category	25 to 34	351	30.1%
			35 to 44	370	31.8%
			45 to 54	260	22.3%
			55 to 64	184	15.8%
			Total	1165	100.0%
	Total	Age category	25 to 34	627	29.8%
			35 to 44	679	32.2%
			45 to 54	481	22.8%
			55 to 64	320	15.2%
			Total	2107	100.0%

Figura 24. Grandi totali per un tavolo nidificato

Si noti che il totale complessivo è di soli 2.107, non 2.828. Due categorie di età sono ancora escluse dalla tabella, quindi i casi in quelle categorie sono esclusi da tutti i totali.

Totale variabili di strato

I totali per le variabili di strato vengono visualizzati come strati separati nella tabella.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Clicca su **Layers** nel builder da tavolo per visualizzare l'elenco Layers.
3. Trascinare e rilasciare *Gender* dall'area di riga sul riquadro canvas alla lista dei Layers.

Nota: da quando hai già specificato i totali per *Gender*, non è necessario farlo ora. Spostando la variabile tra le dimensioni non influiscono nessuna delle impostazioni per quella variabile.

4. Fare clic su **OK** per creare la tabella.
5. Fare doppio clic sulla tabella nel Viewer per attivarla.
6. Clicca sulla freccia in basso nell'elenco a discesa Layer per visualizzare un elenco di tutti gli strati della tabella.

Ci sono tre strati nella tabella: *Genere Maschile*, *Gender Femalee* *Gender Total*.

Visualizza posizione di Layer Totals

Per i totali variabili di strato, la posizione di visualizzazione (sopra o sotto) per i totali determina la posizione di strato per i totali. Ad esempio, se si specificano **Sopra categorie a cui si applicano** per un totale variabile di strato, lo strato totale è il primo strato visualizzato.

Totale parziali

È possibile includere subtotali per sottoinsiemi di categorie di una variabile. Ad esempio, si potrebbero includere i subtotali per le categorie di età che rappresentano tutti gli intervistati nell'indagine campionaria sotto e oltre l'età 45.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Clicca su **Reset** per cancellare le eventuali impostazioni precedenti nel builder table.
3. Nel builder table, trascinare e rilasciare la *Categoria Età* dall'elenco delle variabili all'area Rows sul riquadro canvas.
4. Fare clic con il tasto destro del mouse su *Age category* sul riquadro canvas e scegliere **Categorie e Totali** dal menu a comparsa.
5. Selezionare **3.00** nell'elenco Valori.
6. Clicca su **Aggiungi Subtotale** per visualizzare la finestra di dialogo Definisci Sottototale.
7. Nel campo di testo Etichetta, immettere Subtotal < 45.

8. Quindi fare clic su **Continua**.

Questo inserisce una riga contenente il sottototale per le prime tre categorie di età.

9. Selezionare **6.00** nell'elenco Valori.

10. Clicca su **Aggiungi Sottototale** per visualizzare la finestra di dialogo Definisci Sottototale.

11. Nel campo di testo Etichetta, immettere Subtotal 45+.

12. Quindi fare clic su **Continua**.

Nota importante: è necessario selezionare la posizione di visualizzazione per i totali e i totali parziali (**Sopra categorie a cui si applicano** o **Categorie di Below a cui si applicano**) prima di definire eventuali subtotali. La modifica della posizione del display interessa tutti i subtotali (non solo il sottototale attualmente selezionato) e inoltre *cambia le categorie incluse nei subtotali*.

13. Fare clic su **Applica** e quindi fare clic su **OK** nel builder da tavolo per creare la tabella.

		Count
Age category	Less than 25	242
	25 to 34	627
	35 to 44	679
	Subtotal < 45	1548
	45 to 54	481
	55 to 64	320
	65 or older	479
	Subtotal 45+	1280

Figura 25. Subtotali per la categoria Età

Ciò Che Vedete È Quello Che Viene Subtotalato

Proprio come i totali, i subtotali si basano sulle categorie incluse nella tabella.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).

2. Fare clic con il tasto destro del mouse su *Age category* sul riquadro canvas e scegliere **Categorie e Totali** dal menu a comparsa.

Nota: il valore (non l'etichetta) visualizzato per il primo subtotal è **1.00...3.00**, che indica che il sottototale include tutti i valori nell'elenco compresi tra 1 e 3.

3. Selezionare **1.00** nell'elenco Valori (o fare clic sull'etichetta *Meno di 25*).

4. Fare clic sul tasto freccia a sinistra dell'elenco Escludi.

La prima categoria di età viene ora esclusa e il valore visualizzato per il primo totale parziale cambia in **2.00...3.00**, indicando che la categoria esclusa non verrà inclusa nel totale parziale perché i totali parziali si basano sulle categorie incluse nella tabella. L'esclusione di una categoria lo esclude automaticamente da eventuali subtotali, quindi non è possibile, ad esempio, visualizzare solo i subtotali senza le categorie su cui si basano i subtotali.

Nascondere Categorie subtotali

È possibile sopprimere la visualizzazione delle categorie che definiscono un sottototale e visualizzare solo le categorie subtotali, effettivamente "comprimendo" senza influenzare i dati sottostanti.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).

2. Clicca su **Reset** per cancellare le eventuali impostazioni precedenti nel builder table.

3. Nel builder table, trascinare e rilasciare la *Categoria Età* dall'elenco delle variabili all'area Rows sul riquadro canvas.

4. Fare clic con il tasto destro del mouse su *Age category* sul riquadro canvas e scegliere **Categorie e Totali** dal menu a comparsa.

5. Selezionare **3.00** nell'elenco Valori.

6. Clicca su **Aggiungi Sottototale** per visualizzare la finestra di dialogo Definisci Sottototale.

7. Nel campo di testo Etichetta, immettere Less than 45.
8. Selezionare (controllare) **Nascondi categorie sottototali dalla tabella**.
9. Quindi fare clic su **Continua**.

Questo inserisce una riga contenente il sottototale per le prime tre categorie di età.

10. Selezionare **6.00** nell'elenco Valori.
11. Clicca su **Aggiungi Sottototale** per visualizzare la finestra di dialogo Definisci Sottototale.
12. Nel campo di testo Etichetta, immettere 45 or older.
13. Selezionare (controllare) **Nascondi categorie sottototali**.
14. Quindi fare clic su **Continua**.
15. Per includere un totale con i subtotali, selezionare (controllare) **Totale** nel gruppo Mostra.
16. Fare clic su **Applica**.

La canvas riflette il fatto che verranno visualizzati i subtotali ma saranno escluse le categorie che definiscono i subtotali.

17. Clicca su **OK** per produrre la tabella.

		Count
Age category	Less than 45	1548
	45 or older	1280
	Total	2828

Figura 26. Tabella visualizzazione solo subtotali e totali

Sottototali variabili di strato

Esattamente come i totali, i subtotali per le variabili di strato vengono visualizzati come strati separati nella tabella. Essenzialmente, i subtotali sono trattati come categorie. Ogni categoria è uno strato separato nella tabella e l'ordine di visualizzazione delle categorie di strati è determinato dall'ordine di categoria specificato nella finestra di dialogo Categorie e Totali, inclusa la posizione di visualizzazione delle categorie sottototali.

Categorie Calcolate per Variabili categoriali

È possibile includere categorie calcolate in tabelle personalizzate. Si tratta di nuove categorie che vengono calcolate da categorie della stessa variabile a qualsiasi livello di nidificazione in qualsiasi dimensione - riga, colonna o strato. Ad esempio, si potrebbe includere una categoria calcolata che mostra la differenza tra due categorie.

File di dati di esempio

Gli esempi in questo capitolo utilizzano il file di dati *survey_sample.sav*. Per ulteriori informazioni, consultare l'argomento ["File di esempio"](#) a pagina 76.

Categoria semplice Computed

1. Dai menu, scegliere:
 - Analisi > Tavole > Tabelle personalizzate ...**
2. Nel builder table, trascinare e rilasciare la *Categoria Età* dall'elenco delle variabili all'area Rows sul riquadro canvas.
3. Fare clic con il tasto destro del mouse su *Age category* sul riquadro canvas e scegliere **Categorie e Totali** dal menu a comparsa.
4. Selezionare **3.00** nell'elenco Valori.
5. Fare clic su **Aggiungi categoria** per visualizzare la finestra di dialogo Definisci categoria Compute.
6. Nel campo di testo Etichetta per categoria calcolata, immettere Less than 45.

7. Selezionare **Meno di 25 (1.00)** nell'elenco Categorie e fare clic sulla freccia per copiarlo nella casella di testo Espressione per categoria calcolata. [1] viene visualizzato nell'espressione.
8. Fare clic sul pulsante dell'operatore più (+) nella finestra di dialogo (o premere il tasto + sulla tastiera).
9. Selezionare **Da 25 a 34 (2.00)** nell'elenco Categorie e fare clic sul tasto freccia per copiarlo nella casella di testo Espressione per categoria calcolata.
10. Fare clic sul pulsante dell'operatore più (+) nella finestra di dialogo (o premere il tasto + sulla tastiera).
11. Selezionare **da 35 a 44 (3.00)** nell'elenco Categorie e fare clic sulla freccia per copiarlo nella casella di testo Espressione per categoria calcolata.
12. Quindi fare clic su **Continua**.

Questo inserisce una riga contenente il sottototale per le prime tre categorie di età.

13. Selezionare **5.00** nell'elenco Valori.
14. Clicca su **Aggiungi Sottototale** per visualizzare la finestra di dialogo Definisci Sottototale.
15. Nel campo di testo Etichetta, immettere Less than 65.
16. Quindi fare clic su **Continua**.

Questo inserisce una riga contenente il sottototale per le prime cinque categorie.

17. Fare clic su **Applica** e quindi fare clic su **OK** nel builder da tavolo per creare la tabella.

		Count
Age category	Less than 25	242
	25 to 34	627
	35 to 44	679
	Less than 45	1548
	45 to 54	481
	55 to 64	320
	Less than 65	2349
	65 or older	479

Figura 27. Categoria calcolata con sottototale

La tabella include una categoria calcolata (*Meno di 45*) e un sottototale (*Meno di 65*). Il sottototale include categorie incluse anche nella categoria computata. Non è stato possibile creare la stessa tabella con i soli sottotali, perché i sottotali non possono condividere le stesse categorie.

Nascondere Categorie in una categoria Computed

Come per i sottotali, è possibile sopprimere la visualizzazione delle categorie che vengono utilizzate in un'espressione di categoria computata e visualizzare solo la categoria computata stessa. Il seguente esempio si costruisce sul precedente.

1. Dai menu, scegliere:
Analisi > Tavole > Tabelle personalizzate ...
2. Fare clic con il tasto destro del mouse su *Age category* sul riquadro canvas e scegliere **Categorie e Totali** dal menu a comparsa.
3. Selezionare la categoria computata *Meno di 45* nella lista Value (s).
4. Cliccare su **Modifica** per visualizzare la finestra di dialogo Definisci categoria Compute.
5. Selezionare **Categorie Nascondi utilizzate in espressione dalla tabella**.
6. Quindi fare clic su **Continua**.
7. Selezionare il totale *Meno di 65* nell'elenco Value (s).
8. Clicca su **Modifica** per visualizzare la finestra di dialogo Definisci Sottototale.
9. Selezionare **Nascondi categorie sottotali dalla tabella**.

10. Quindi fare clic su **Continua**.

11. Fare clic su **Applica** e quindi fare clic su **OK** nel builder da tavolo per creare la tabella.

		Count
Age category	Less than 45	1548
	Less than 65	2349
	65 or older	479

Figura 28. Categoria computata con categorie subtotali e nascoste

Come l'esempio precedente, la tabella include una categoria computata e un sottototale. Ma in questo caso le categorie in ognuno sono nascoste in modo che vengano mostrati solo questi totali.

Subtotali di riferimento in una categoria Computed

È possibile includere i subtotali in un'espressione di categoria computata.

1. Dai menu, scegliere:

Analisi > Tavole > Tabelle personalizzate ...

2. Clicca su **Reset** per cancellare le eventuali impostazioni precedenti nel builder table.

3. Nel builder table, trascinare e rilasciare *Stato forza lavoro* dall'elenco di variabili nell'area Rows del riquadro canvas.

4. Trascinare e rilasciare *stato civile* dall'elenco delle variabili nell'area Colonne.

5. Fare clic con il tasto destro del mouse su *Stato forza lavoro* sul riquadro canvas e scegliere **Categorie e Totali** dal menu a comparsa.

6. Selezionare **2** nella lista Value (s).

7. Clicca su **Aggiungi Sottototale** per visualizzare la finestra di dialogo Definisci Sottototale.

8. Nel campo di testo Etichetta, immettere *Working*.

9. Selezionare **Nascondi categorie sottototali dalla tabella**.

10. Quindi fare clic su **Continua**.

Questo inserisce una riga contenente il sottototale per le prime due categorie di stato di lavoro.

11. Selezionare **8** nella lista Value (s).

12. Clicca su **Aggiungi Sottototale** per visualizzare la finestra di dialogo Definisci Sottototale.

13. Nel campo di testo Etichetta, immettere *Not Working*.

14. Selezionare **Nascondi categorie sottototali**.

15. Quindi fare clic su **Continua**.

Questo inserisce una riga contenente il sottototale per le altre categorie di stato di lavoro.

16. Selezionare il sottototale di *Non di lavoro* nella lista Value (s).

17. Fare clic su **Aggiungi categoria** per visualizzare la finestra di dialogo Definisci categoria Compute.

18. Nel campo di testo Etichetta per categoria calcolata, immettere *Working / Not Working*.

19. Selezionare **Lavorare (Working #1)** nella lista Totali e Subtotali e fare clic sul pulsante freccia per copiarlo nella casella di testo Expression for Computed Category.

20. Fare clic sul pulsante di operatore (/) operatore nella finestra di dialogo (o premere il / tasto sulla tastiera).

21. Selezionare **Non Lavorare (Non di lavoro #2)** nella lista Totali e Subtotali e fare clic sul pulsante freccia per copiarlo nella casella di testo Expression for Computed Category.

Per impostazione predefinita, la categoria computata utilizza lo stesso formato della statistica della variabile, che è Conte in questo caso. Poiché vogliamo mostrare posizioni decimali derivanti dalla divisione nell'espressione della categoria computata e il formato predefinito per il conteggio non include posizioni decimali, è necessario modificare il formato.

22. Fare clic sulla scheda Visualizza formati.
23. Modificare l'impostazione Decimals per il conteggio a **2**.
24. Quindi fare clic su **Continua**.
25. Fare clic su **Applica** e quindi fare clic su **OK** nel builder da tavolo per creare la tabella.

		Marital status				
		Married	Widowed	Divorced	Separated	Never married
		Count	Count	Count	Count	Count
Labor force status	Working	916	64	330	67	494
	Not Working	429	219	116	26	169
	Working / Not Working	2.14	.29	2.84	2.58	2.92

Figura 29. Categoria di calcolo che mostra il rapporto dei subtotali

La tabella include due totali parziali e una categoria computata. La categoria computata mostra il rapporto dei subtotali in modo da poter facilmente confrontare i gruppi rappresentati da ciascun sottotale. C'è un rapporto molto più basso di lavoro per non lavorare gli intervistati vedenti rispetto agli altri gruppi. Inoltre, c'è un rapporto leggermente più basso tra gli intervistati sposati, magari derivanti da coniugi che lasciano la forza lavoro per rimanere a casa con un bambino.

Utilizzo Categorie Calcolate per visualizzare Subtotali non esaustivi

I subtotali sono esaustivi. Ovvero tutti i subtotali in una tabella includono tutti i valori sopra o sotto le loro posizioni in tavola. Le categorie calcolate, invece, non sono esaustive e consentono di riassumere un mix di categorie in una tabella.

1. Dai menu, scegliere:

Analisi > Tavole > Tabelle personalizzate ...

2. Clicca su **Reset** per cancellare le eventuali impostazioni precedenti nel builder table.
3. Nel builder table, trascinare e rilasciare *Think of self as liberal o conservatore* dall'elenco delle variabili nell'area Rows del riquadro canvas.
4. Fare clic con il tasto destro del mouse su *Think di sé come liberal o conservativo* sul riquadro canvas e scegliere **Categorie e Totali** dal menu del contesto pop - up.
5. Selezionare **3** nella lista Value (s).
6. Clicca su **Aggiungi categoria** per visualizzare la finestra di dialogo Definisci categoria Computed.
7. Nel campo di testo Etichetta per categoria calcolata, immettere **Liberal Subtotal**. Da notare che ci sono quattro spazi prima del testo. Questi spazi sono utilizzati per l'indentazione nella tabella risultante.
8. Selezionare **Estremamente liberale (1)** nell'elenco Categorie e fare clic sul pulsante freccia per copiarlo nella casella di testo Expression for Computed Category.
9. Fare clic sul pulsante dell'operatore più (+) nella finestra di dialogo (o premere il tasto + sulla tastiera).
10. Selezionare **Liberale (2)** nell'elenco Categorie e fare clic sul pulsante freccia per copiarlo nella casella di testo Expression for Computed Category.
11. Fare clic sul pulsante dell'operatore più (+) nella finestra di dialogo (o premere il tasto + sulla tastiera).
12. Selezionare **Sleggereliberale (3)** nell'elenco Categorie e fare clic sul pulsante freccia per copiarlo nella casella di testo Expression for Computed Category.
13. Fare clic su **Continua**.

Questo inserisce una riga contenente il sottototale per le categorie liberali.

14. Selezionare **7** nella lista Value (s).
15. Clicca su **Aggiungi categoria** per visualizzare la finestra di dialogo Definisci categoria Computed.

16. Nel campo di testo Etichetta per categoria calcolata, immettere **Conservative Subtotal**. Da notare che ci sono quattro spazi prima del testo. Questi spazi sono utilizzati per l'indentazione nella tabella risultante.
17. Selezionare **Slight conservative (5)** nell'elenco Categorie e fare clic sul pulsante freccia per copiarlo nella casella di testo Expression for Computed Category.
18. Fare clic sul pulsante dell'operatore più (+) nella finestra di dialogo (o premere il tasto + sulla tastiera).
19. Selezionare **Conservatore (6)** nell'elenco Categorie e fare clic sul pulsante freccia per copiarlo nella casella di testo Expression for Computed Category.
20. Fare clic sul pulsante dell'operatore più (+) nella finestra di dialogo (o premere il tasto + sulla tastiera).
21. Selezionare **Estremamente conservativo (7)** nell'elenco Categorie e fare clic sul pulsante freccia per copiarlo nella casella di testo Expression for Computed Category.
22. Fare clic su **Continua**.

Questo inserisce una riga contenente il sottototale per le categorie conservatrici.

23. Fare clic su **Applica** e quindi fare clic su **OK** nel builder da tavolo per creare la tabella.

	Count
Think of self as liberal or conservative	
Extremely liberal	64
Liberal	357
Slightly liberal	351
Liberal Subtotal	772
Moderate	986
Slightly conservative	432
Conservative	415
Extremely conservative	86
Conservative Subtotal	933

Figura 30. Categorie calcolate che visualizzano subtotali non esaustivi

La tabella include due categorie calcolate che non includono tutte le categorie visualizzate nella tabella. La categoria *Moderato* non è inclusa nella categoria computata. Non è possibile creare la stessa tabella con i subtotali perché i subtotali sono esaustivi.

Tabelle per Variabili con Categorie condivise

Le indagini spesso contengono molte domande con un insieme comune di possibili risposte. Ad esempio, il nostro sondaggio campione contiene una serie di variabili riguardanti la fiducia in varie istituzioni e servizi pubblici e privati, tutti con la stessa serie di categorie di risposta: 1 = *Un grande affare*, 2 = *Solo alcune* 3 = *Hardly any*. È possibile utilizzare la sovrapposizione per visualizzare queste variabili correlate nella stessa tabella -- ed è possibile visualizzare le categorie di risposta condivise nelle colonne della tabella. Queste funzioni sono disponibili anche se si utilizzano categorie calcolate, con la disposizione che qualsiasi etichetta e espressione della categoria computata sono le stesse in tutte le variabili.

	A great deal	Only some	Hardly any
Confidence in banks & financial institutions	490	1068	306
Confidence in education	511	1055	315
Confidence in major companies	500	1078	243
Confidence in medicine	844	864	167
Confidence in press	176	878	808
Confidence in television	196	936	744

Figura 31. Tabella delle variabili con categorie condivise

Nota: nella versione precedente delle Tabelle personalizzate, questa era conosciuta come "tabella delle frequenze".

File di dati di esempio

Gli esempi in questo capitolo utilizzano il file di dati *survey_sample.sav*. Per ulteriori informazioni, consultare l'argomento "File di esempio" a pagina 76.

Tutti gli esempi forniti qui visualizzano etichette variabili in finestre di dialogo, ordinate in ordine alfabetico. Le proprietà di visualizzazione della lista variabili sono impostate sulla scheda Generale nella finestra di dialogo Opzioni (Modifica menu, Opzioni).

Tabella dei Conti

1. Dai menu, scegliere:

Analisi > Tavole > Tabelle personalizzate ...

2. Nell'elenco delle variabili nel builder table, clicca su *Confidence nelle banche ...* e poi Shift - clicca su *Confidence in televisione* per selezionare tutte le variabili "di fiducia". (*Nota:* questo presuppone che le etichette variabili siano visualizzate in ordine alfabetico, non di file, nell'elenco delle variabili.)
3. Trascinare e rilasciare le sei variabili di confidenza all'area Rows sul riquadro canvas.

Questo accelera le variabili nella dimensione della riga. Per impostazione predefinita, le etichette di categoria per ogni variabile vengono visualizzate anche nelle righe, risultando in una tabella molto lunga, ristretta (6 variabili x 3 categorie = 18 righe) -- ma dal momento che tutte e sei le variabili condividono le stesse etichette di categoria definite (etichette di valore), è possibile inserire le etichette di categoria nella dimensione della colonna.

4. Dall'elenco a discesa Posizione Categoria, selezionare **Label di riga in Colonne**.

Ora la tabella ha solo sei righe, una per ciascuna delle variabili in pila e le categorie definite diventano colonne nella tabella.

5. Prima di creare la tabella, selezionare (clicca) **Nascondi** per Posizione nel gruppo Statistiche di riepilogo, dal momento che l'etichetta statistica di riepilogo *Conteggio* non è davvero necessaria.
6. Fare clic su **OK** per creare la tabella.

	A great deal	Only some	Hardly any
Confidence in banks & financial institutions	490	1068	306
Confidence in education	511	1055	315
Confidence in major companies	500	1078	243
Confidence in medicine	844	864	167
Confidence in press	176	878	808
Confidence in television	196	936	744

Figura 32. Tabella delle variabili di riga in sovrapposizione con etichette di categoria condivise in colonne

Invece di visualizzare le variabili nelle righe e nelle categorie nelle colonne, è possibile creare una tabella con le variabili impilate nelle colonne e le categorie visualizzate nelle righe. Questa potrebbe essere una scelta migliore se ci fossero più categorie che variabili, mentre nel nostro esempio ci sono più variabili che categorie.

Tabella delle percentuali

Per una tabella con variabili sovrapposte a righe e categorie visualizzate in colonne, la percentuale più significativa (o almeno più semplice da capire) da visualizzare sono le percentuali di riga. (Per una tabella con variabili incastonate nelle colonne e nelle categorie visualizzate nelle righe, si vorrebbero probabilmente le percentuali di colonna.)

1. Aprire il builder table di nuovo (Analizza menu, tabelle, Tabelle personalizzate).

2. Fare clic con il tasto destro del mouse su una qualsiasi delle variabili di confidenza presenti nell'anteprima della tabella sul riquadro canvas e scegliere **Statistiche di riepilogo** dal menu a comparsa.
3. Selezionare **Riga N%** nell'elenco Statistiche e fare clic sul pulsante freccia per spostarlo nella lista Visualizza.
4. Clicca su qualsiasi cella nella riga *Conteggio* nell'elenco Visualizza e clicca sul pulsante freccia per spostarlo nuovamente nell'elenco Statistiche, rimuoverlo dall'elenco Visualizza.
5. Clicca su **Applica a tutto** per applicare la modifica statistica di riepilogo a tutte le variabili impilate nella tabella.

Nota: se la tua anteprima della tabella non sembra questa figura, probabilmente clicca **Applica alla selezione** invece di **Applica a tutto**, che applica la nuova statistica di riepilogo solo alla variabile selezionata. In questo esempio, ciò si tradurrebbe in due colonne per ogni categoria: una con segnaposto visualizzati per tutte le altre variabili e una con segnaposto percentuale riga visualizzata per la variabile selezionata. Questo è esattamente il tavolo che sarebbe prodotto ma *non* quello che vogliamo in questo esempio.

6. Fare clic su **OK** per creare la tabella.

	A great deal	Only some	Hardly any
Confidence in banks & financial institutions	26.3%	57.3%	16.4%
Confidence in education	27.2%	56.1%	16.7%
Confidence in major companies	27.5%	59.2%	13.3%
Confidence in medicine	45.0%	46.1%	8.9%
Confidence in press	9.5%	47.2%	43.4%
Confidence in television	10.4%	49.9%	39.7%

Figura 33. Tabella delle percentuali di riga per le variabili incastonate in righe, categorie visualizzate in colonne

Nota: è possibile includere qualsiasi numero di statistiche di riepilogo in una tabella di variabili con categorie condivise. I nostri esempi mostrano solo uno alla volta per mantenerli semplici.

Totali e controllo di categoria

È possibile creare tabelle con categorie nella dimensione opposta dalle variabili solo se tutte le variabili presenti nella tabella hanno le stesse categorie, visualizzate nello stesso ordine. Questo include totali, subtotali e qualsiasi altra regolazioni di categoria che effettui. Ciò significa che le eventuali modifiche apportate nella finestra di dialogo Categorie e Totali devono essere effettuate per tutte le variabili presenti nella tabella che condividono le categorie.

1. Aprire il builder table di nuovo (Analizza menu, tabelle, Tabelle personalizzate).
2. Fare clic con il tasto destro del mouse sulla prima variabile di confidenza nella tabella in anteprima sul riquadro canvas e scegliere **Categorie e Totali** dal menu a comparsa.
3. Selezionare (controllare) **Totale** nella finestra di dialogo Categorie e Totali e quindi fare clic su **Applica**.

La prima cosa che probabilmente noterete è che le etichette di categoria si sono spostate dalle colonne alle righe. Si può anche notare che il controllo della posizione di categoria è ora disabilitato. Questo perché le variabili non condividono più lo stesso identico insieme di "categorie". Una delle variabili ora ha una categoria totale.

4. Fare clic con il tasto destro del mouse su una qualsiasi delle variabili di confidenza sul riquadro canvas e selezionare **Seleziona tutte le variabili di riga** dal menu a comparsa - oppure Ctrl - clicca su ogni variabile impilata sul riquadro canvas fino a quando non sono tutti selezionati (è possibile che sia necessario scorrere il riquadro o espandere la finestra del builder della tabella).
5. Fare clic su **Categorie e totali** nel gruppo Definisci.

6. Se **Total** non è già selezionato (controllato) nella finestra di dialogo Categorie e Totali, selezionarlo ora e quindi fare clic su **Applica**.
7. L'elenco a discesa Posizione Categoria deve essere nuovamente abilitato, dato che ora tutte le variabili hanno la categoria totale aggiuntiva, quindi selezionare **Label Row in Colonne**.
8. Fare clic su **OK** per creare la tabella.

	A great deal	Only some	Hardly any	Total
Confidence in banks & financial institutions	26.3%	57.3%	16.4%	100.0%
Confidence in education	27.2%	56.1%	16.7%	100.0%
Confidence in major companies	27.5%	59.2%	13.3%	100.0%
Confidence in medicine	45.0%	46.1%	8.9%	100.0%
Confidence in press	9.5%	47.2%	43.4%	100.0%
Confidence in television	10.4%	49.9%	39.7%	100.0%

Figura 34. Tabella delle percentuali di riga per le variabili impilate in righe, categorie e totali visualizzati nelle colonne

Nidificazione in tabelle con Categorie condivise

Nelle tabelle nidificate, le variabili impilate con le categorie condivise devono essere al livello più intime della loro dimensione se si desidera visualizzare le etichette di categoria nella dimensione opposta.

1. Aprire il builder table di nuovo (Analizza menu, tabelle, Tabelle personalizzate).
2. Trascinare e rilasciare *Gender* dall'elenco delle variabili al lato sinistro dell'area Rows.

Le variabili impilate con categorie condivise sono ora nidificate all'interno delle categorie di genere nell'anteprima della tabella.

3. Ora trascinare e rilasciare *Gender* a destra di una delle variabili di confidenza impilate nell'anteprima della tabella.

Ancora una volta, le etichette di categoria si sono invertite alla quota di fila, e il controllo della posizione di categoria è disabilitato. Si dispone ora di una variabile impilata che ha anche *Gender* nidificata al suo interno, mentre le altre variabili in sovrapposizione non contengono variabili nidificate. Si potrebbe aggiungere *Gender* come variabile nidificata a ciascuna delle variabili in sovrapposizione, ma poi spostare le etichette di riga su colonne risulterebbe nelle etichette di categoria per *Gender* visualizzate nelle colonne, non le etichette di categoria per le variabili in sovrapposizione con le categorie condivise. Questo perché *Gender* sarebbe ora la variabile innerpiù nidificata e la modifica della posizione di categoria si applica sempre alla variabile innerpiù nidificata.

Statistiche di riepilogo

Statistiche di riepilogo includono tutto da semplici conteggi per variabili categoriali a misure di dispersione, come l'errore standard della media per le variabili di scala. Esso *non* include i test di significatività disponibili nella scheda Statistiche di test nella finestra di dialogo Tabelle personalizzate. Per ulteriori informazioni, consultare l'argomento [“Statistiche del test” a pagina 53](#).

Le statistiche di riepilogo per le variabili categoriali e i set di risposta multipla includono conteggi e un'ampia varietà di calcoli percentuali, tra cui:

- Percentuali riga
- Percentuali colonna
- Percentuali sottotabella
- Percentuali di tabella
- Percentuali N valide

Oltre alle statistiche di riepilogo disponibili per le variabili categoriali, le statistiche di riepilogo per le variabili di scala e i riepiloghi totali personalizzati per le variabili categoriali includono:

- Media
- Mediana
- Percentili
- Somma
- Deviazione standard
- Intervallo
- valori minimi e massimi

Ulteriori statistiche di riepilogo sono disponibili per più set di risposta. È disponibile anche un elenco completo di statistiche di riepilogo. Per ulteriori informazioni, consultare l'argomento [“Statistiche di riepilogo”](#) a pagina 6.

File di dati di esempio

Gli esempi in questo capitolo utilizzano il file di dati *survey_sample.sav*. Per ulteriori informazioni, consultare l'argomento [“File di esempio”](#) a pagina 76.

Tutti gli esempi forniti qui visualizzano etichette variabili in finestre di dialogo, ordinate in ordine alfabetico. Le proprietà di visualizzazione della lista variabili sono impostate sulla scheda Generale nella finestra di dialogo Opzioni (Modifica menu, Opzioni).

Variabile di origine statistiche di riepilogo

Le statistiche di riepilogo disponibili dipendono dal livello di misurazione della variabile di origine delle statistiche di riepilogo. La fonte delle statistiche di riepilogo (la variabile su cui si basano le statistiche di riepilogo) è determinata da:

- **Livello di misurazione.** Se una tabella (o una sezione di tabella in una tabella in sovrapposizione) contiene una variabile di scala, le statistiche di riepilogo si basano sulla variabile di scala.
- **Ordine di selezione variabile.** La dimensione di origine delle statistiche predefinite (riga o colonna) per le variabili categoriali si basa sull'ordine in cui si trascina e si abbassano le variabili sul riquadro canvas. Ad esempio, se si trascina prima una variabile all'area delle righe, la dimensione della riga è la dimensione di origine delle statistiche predefinite.
- **Nidificazione.** Per le variabili categoriali, le statistiche di riepilogo si basano sulla variabile innerpiù nella dimensione di origine delle statistiche.

Una tabella in sovrapposizione può avere più variabili di origine statistiche di riepilogo (sia scala che categoriale), ma ogni sezione di tabella ha una sola fonte di statistiche di riepilogo.

Riepilogo statistiche Fonte per Variabili categoriali

1. Dai menu, scegliere:

Analisi > Tavole > Tabelle personalizzate ...

2. Nel builder table, trascinare e rilasciare la *Categoria Età* dall'elenco delle variabili nell'area Rows del riquadro canvas.
3. Fare clic con il tasto destro del mouse su *Age category* sul riquadro canvas e selezionare **Statistiche di riepilogo** dal menu a comparsa. (Poiché questa è l'unica variabile nella tabella, è la variabile di origine statistica.)
4. Nella finestra di dialogo Statistiche di riepilogo, selezionare *Colonna N%* nell'elenco Statistiche e fare clic sulla freccia per aggiungendola alla lista Visualizza.
5. Fare clic su **Applica alla selezione**.
6. Nel builder table, trascinare e rilasciare *Ottieni notizie da internet* a destra della *Categoria Età* sul riquadro canvas.

7. Fare clic con il tasto destro del mouse su *Age category* sul riquadro canvas. L'articolo **Statistiche di riepilogo** sul menu a comparsa è ora disabilitato in quanto *Categoria Età* non è la variabile innerpiù nidificata nella dimensione di origine delle statistiche.
8. Fare clic con il tasto destro del mouse su *Ottieni notizie da internet* sul riquadro canvas. L'articolo **Statistiche di riepilogo** è abilitato perché è ora la variabile di origine delle statistiche di riepilogo, dato che è la variabile innerpiù nidificata nella dimensione di origine delle statistiche. (Dal momento che la tabella ha una sola dimensione - righe - è la dimensione della fonte di statistica.)
9. Trascinare e rilasciare *Ottieni notizie da internet* dall'area Rows sul riquadro canvas nell'area delle Colonne.
10. Fare clic con il tasto destro del mouse su *Ottieni notizie da internet* sul riquadro canvas. L'articolo **Statistiche di riepilogo** sul menu a comparsa è ora disabilitato perché la variabile non è più nella dimensione di origine delle statistiche.

La categoria *Age* è ancora una volta la variabile di origine statistica perché la dimensione di origine delle statistiche predefinite per le variabili categoriali è la prima dimensione in cui si mettono le variabili quando si crea la tabella. In questo esempio, la prima cosa che abbiamo fatto è stata mettere delle variabili nella dimensione della riga. Così, la dimensione della riga è la dimensione di origine delle statistiche di default; e siccome la *Categoria Età* è ora l'unica variabile in quella dimensione, è la variabile di origine statistica.

Statistiche di riepilogo Origine per le variabili di scala

1. Trascinare e rilasciare la variabile di scala *Ore al giorno guardando la TV* a sinistra della *Categoria Età* nell'area Rows del riquadro canvas.
- La prima cosa che si potrebbe notare è che i riepiloghi *Conteggio* e *Colonna N%* sono stati sostituiti con *Mean--* e se si fa clic con il tasto destro del mouse su *Ore al giorno guardando la TV* sul riquadro canvas, vedrete che ora è la variabile di origine delle statistiche di riepilogo. Per una tabella con una variabile di scala, la variabile di scala è sempre la variabile di origine statistica indipendentemente dal suo livello di nidificazione o dimensione, e la statistica di riepilogo predefinita per le variabili di scala è la media.
2. Trascinare e rilasciare *Ore al giorno guardando la TV* dall'area Rows nell'area delle Colonne sopra *Ottieni notizie da internet*.
3. Fare clic con il tasto destro del mouse su *Ore al giorno guardando la TV* e selezionare **Statistiche di riepilogo** dal menu a comparsa. (E'ancora la variabile di origine statistica anche quando lo si sposta in una dimensione diversa.)
4. Nella finestra di dialogo Statistiche di riepilogo, fare clic sulla cella **Formato** per la media nell'elenco Visualizza e selezionare **nnnn** dall'elenco a discesa Formato. (Potrebbe essere necessario scorrere l'elenco per trovare questa scelta.)
5. Nella cella Decimali, immettere 2.
6. Fare clic su **Applica alla selezione**.

L'anteprima della tabella sul riquadro canvas mostra ora che i valori medi verranno visualizzati con due decimali.

7. Fare clic su **OK** per creare la tabella.

		Hours per day watching TV	
		Get news from internet	
		No	Yes
		Mean	Mean
Age category	Less than 25	3.54	2.12
	25 to 34	3.42	2.14
	35 to 44	3.00	2.01
	45 to 54	2.83	2.06
	55 to 64	3.24	2.37
	65 or older	3.82	2.33

Figura 35. Variabile di scala riepilogata all'interno di variabili categoriali crosstabulate

Variabili Stack

Dal momento che una tabella in sovrapposizione può contenere più variabili di origine statistiche e si possono specificare statistiche di riepilogo diverse per ognuna di quelle variabili di origine statistiche, ci sono alcune considerazioni speciali per specificare le statistiche di riepilogo in tabelle impilate.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Clicca su **Reset** per cancellare le eventuali impostazioni precedenti nel builder table.
3. Clicca su *Ottieni notizie da internet* nell'elenco delle variabili e poi fai clic su *Ottieni notizie dalla televisione* nell'elenco delle variabili per selezionare tutte le variabili "news". (Nota: questo presuppone che le etichette variabili siano visualizzate in ordine alfabetico, non di file, nell'elenco delle variabili.)
4. Trascinare e rilasciare le cinque variabili di notizie nell'area Rows del riquadro canvas.
Le cinque variabili di notizie sono impilate nella dimensione della riga.
5. Clicca su *Ottieni notizie da internet* sul riquadro canvas in modo che sia selezionata solo quella variabile.
6. Ora clicca con il tasto destro del mouse su *Ottieni notizie da internet* e seleziona **Statistiche riepilogo** dal menu a comparsa.
7. Nella finestra di dialogo Statistiche di riepilogo, selezionare *Colonna N%* dall'elenco Statistiche e fare clic sulla freccia per aggiungendola alla lista Visualizza. (È possibile utilizzare la freccia per spostare le statistiche selezionate dall'elenco Statistiche nell'elenco Visualizza oppure si possono trascinare e rilasciare statistiche selezionate dall'elenco Statistiche nell'elenco Visualizza.)
8. Quindi fare clic su **Applica alla selezione**.
Viene aggiunta una colonna per le percentuali di colonna -- ma l'anteprima della tabella sul riquadro canvas indica che le percentuali di colonna verranno visualizzate per una sola variabile. Questo perché in una tabella in sovrapposizione ci sono più variabili di origine statistiche e ognuna può avere statistiche di riepilogo diverse. In questo esempio, però, vogliamo visualizzare le stesse statistiche di riepilogo per tutte le variabili.
9. Fare clic con il tasto destro del mouse su *Ottieni notizie dai giornali* sul riquadro canvas e selezionare **Statistiche di riepilogo** dal menu a comparsa.
10. Nella finestra di dialogo Statistiche di riepilogo, selezionare *Colonna N%* dall'elenco Statistiche e fare clic sulla freccia per aggiungendola alla lista Visualizza.
11. Quindi fare clic su **Applica a tutti**.

Ora l'anteprima della tabella indica che le percentuali di colonna verranno visualizzate per tutte le variabili impilate.

Statistiche totali di riepilogo personalizzate per Variabili categoriali

Per le variabili di origine delle statistiche categoriali, è possibile includere statistiche di riepilogo totale personalizzate diverse dalle statistiche visualizzate per le categorie della variabile. Ad esempio, per una variabile ordinale, si potrebbero visualizzare le percentuali per ogni categoria e la media o mediana per la statistica di riepilogo totale personalizzata.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Clicca su **Reset** per cancellare le eventuali impostazioni precedenti nel builder table.
3. Clicca *Confidenza in stampa* nell'elenco delle variabili e quindi Ctrl - clicca su *Confidence in TV* per selezionare entrambe le variabili.
4. Trascinare e rilasciare le due variabili nell'area Rows del riquadro canvas. Questo accelera le due variabili nella dimensione della riga.
5. Fare clic con il tasto destro del mouse sul riquadro canvas e selezionare **Seleziona tutte le variabili di riga** dal menu a comparsa. (Possono entrambi già essere selezionati, ma vogliamo esserne certi.)

6. Fare clic con il tasto destro del mouse sulla variabile e selezionare **Categorie e totali** dal menu a comparsa.
7. Nella finestra di dialogo Categorie e Totali, fare clic (controllare) **Totalee** quindi fare clic su **Applica**.
L'anteprima della tabella sul riquadro canvas ora visualizza una riga totale per entrambe le variabili. Al fine di visualizzare statistiche di riepilogo totale personalizzate, per la tabella devono essere specificati totali e / o subtotali.
8. Fare clic con il tasto destro del mouse sul riquadro canvas e selezionare **Statistiche riepilogo** dal menu a comparsa.
9. Nella finestra di dialogo Statistiche di riepilogo, clicca su *Conteggio* nell'elenco Visualizza e clicca sulla freccia per spostarlo nell'elenco Statistiche, rimuoverlo dall'elenco Visualizza.
10. Clicca su *Colonna N%* nell'elenco Statistiche e clicca sul tasto freccia per spostarlo nella lista Visualizza.
11. Clicca (spunta) **Statistiche riepilogo personalizzate per Totali e Subtotali**.
12. Clicca *Conteggio* nell'elenco Visualizza riepilogo personalizzato e clicca sulla freccia per spostarlo nell'elenco Statistiche di riepilogo personalizzate, rimuovendolo dall'elenco Visualizza.
13. Clicca su *Mean* nell'elenco Statistiche di riepilogo personalizzate e clicca sulla freccia per spostarla nella lista dei Display di riepilogo personalizzati.
14. Fare clic sulla cella **Formato** per la media nell'elenco Visualizza e selezionare **nnnn** dall'elenco a discesa dei formati. (Potrebbe essere necessario scorrere l'elenco per trovare questa scelta.)
15. Nella cella Decimali, immettere 2.
16. Clicca su **Applica ad All** per applicare queste impostazioni a entrambe le variabili nella tabella.
È stata aggiunta una nuova colonna per la statistica di riepilogo totale personalizzata, che potrebbe non essere ciò che si desidera, dato che l'anteprima sul riquadro canvas indica chiaramente che questo si tradurrà in un tavolo con molte celle vuote.
17. Nel builder table, nel gruppo Riepilogo statistiche, selezionare **Righe** dall'elenco a discesa Posizione.
Questo sposta tutte le statistiche di riepilogo alla dimensione della riga, visualizzando tutte le statistiche di riepilogo in una singola colonna della tabella.
18. Fare clic su **OK** per creare la tabella.

Confidence in press	A great deal	Column N %	9.5%
	Only some	Column N %	47.2%
	Hardly any	Column N %	43.4%
	Total	Mean	2.34
Confidence in television	A great deal	Column N %	10.4%
	Only some	Column N %	49.9%
	Hardly any	Column N %	39.7%
	Total	Mean	2.29

Figura 36. Variabili categoriali con statistiche di riepilogo totale personalizzate

Visualizzazione dei valori di categoria

C'è solo un piccolo problema con la tabella precedente - potrebbe essere difficile interpretare il valore medio senza conoscere i valori di categoria sottostanti su cui si basa. Una media di 2.34 è compresa tra *A deal* e *Only some*-- oppure è compresa tra *Only some* e *Hardly any*?

Anche se non possiamo affrontare questo problema direttamente nelle Tabelle personalizzate, possiamo affrontarlo in modo più generale.

1. Dai menu, scegliere:

Modifica > Opzioni...

2. Nella finestra di dialogo Opzioni, fare clic sulla scheda **Label di output**.
3. Nel gruppo Pivot Table Labeling, selezionare **Valori e etichette** dai **Valori variabili nelle etichette mostrate come** elenco a discesa.

4. Clicca su **OK** per salvare questa impostazione.
5. Aprire il builder table (Analyze menu, Tabelle, Tabelle personalizzate) e fare clic su **OK** per creare di nuovo la tabella.

Confidence in press	1 A great deal	Column N %	9.5%
	2 Only some	Column N %	47.2%
	3 Hardly any	Column N %	43.4%
	Total	Mean	2.34
Confidence in television	1 A great deal	Column N %	10.4%
	2 Only some	Column N %	49.9%
	3 Hardly any	Column N %	39.7%
	Total	Mean	2.29

Figura 37. Valori e etichette visualizzati per categorie variabili

I valori di categoria rendono chiaro che una media di 2.34 è compresa tra *Solo alcuni* e *Qualsiasi*. La visualizzazione dei valori di categoria nella tabella rende molto più semplice interpretare il valore delle statistiche di riepilogo totale personalizzate, come la media.

Questa impostazione di visualizzazione è un'impostazione globale che interessa tutte le uscite della tabella pivot da tutte le procedure e persiste nelle sessioni fino a modificarla. Per modificare l'impostazione di ritorno per visualizzare solo etichette di valore:

6. Dai menu, scegliere:
Modifica > Opzioni...
7. Nella finestra di dialogo Opzioni, fare clic sulla scheda **Label di output**.
8. Nel gruppo Pivot Table Labeling, selezionare **Labels** dai **Valori variabili nelle etichette mostrate come** elenco a discesa.
9. Clicca su **OK** per salvare questa impostazione.

Riepilogo Variabili di scala

Riepilogo Variabili di scala

Una vasta gamma di statistiche di riepilogo sono disponibili per le variabili di scala. Oltre ai conteggi e alle percentuali disponibili per le variabili categoriali, le statistiche di riepilogo per le variabili di scala includono anche:

- Media
- Mediana
- Percentili
- Somma
- Deviazione standard
- Intervallo
- valori minimi e massimi

Per ulteriori informazioni, consultare l'argomento [“Statistiche di riepilogo per variabili di scala e Totali personalizzati categoriali” a pagina 9.](#)

File di dati di esempio

Gli esempi in questo capitolo utilizzano il file di dati *survey_sample.sav*. Per ulteriori informazioni, consultare l'argomento [“File di esempio” a pagina 76.](#)

Tutti gli esempi forniti qui visualizzano etichette variabili in finestre di dialogo, ordinate in ordine alfabetico. Le proprietà di visualizzazione di elenco variabili sono specificate nella scheda Generale nella finestra di dialogo Opzioni (Modifica menu, Opzioni).

Variabili Scale in sovrapposizione

È possibile sintetizzare più variabili di scala nella stessa tabella impilandole nella tabella.

1. Dai menu, scegliere:

Analisi > Tavole > Tabelle personalizzate ...

2. Nel builder table, clicca su *Age of respondent* nell'elenco delle variabili, Ctrl - clicca *Highest anno di scuola completato* Ctrl - clicca *Ore al giorno guardando la TV* per selezionare tutte e tre le variabili.
3. Trascinare e rilasciare le tre variabili selezionate all'area Rows del riquadro canvas.

Le tre variabili sono impilate nella dimensione della riga. Poiché tutte e tre le variabili sono variabili di scala, non vengono visualizzate categorie e la statistica di riepilogo predefinita è la media.

4. Fare clic su **OK** per creare la tabella.

	Mean
Age of respondent	46
Highest year of school completed	13
Hours per day watching TV	3

Figura 38. Tabella dei valori medi delle variabili di scala in sovrapposizione

Statistiche di riepilogo multiple

Per impostazione predefinita, la media viene visualizzata per le variabili di scala; tuttavia, è possibile scegliere altre statistiche di riepilogo per le variabili di scala ed è possibile visualizzare più di una statistica di riepilogo.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Fare clic con il tasto destro del mouse su una qualsiasi delle variabili di scala in anteprima della tabella sul riquadro canvas e selezionare **Statistiche di riepilogo** dal menu a comparsa.
3. Nella finestra di dialogo Statistiche di riepilogo, selezionare *Mediana* nell'elenco Statistiche e fare clic sulla freccia per aggiungendola alla lista Visualizza. (è possibile utilizzare la freccia per spostare le statistiche selezionate dall'elenco Statistiche all'elenco Visualizza oppure si possono trascinare e rilasciare statistiche selezionate dall'elenco Statistiche nell'elenco Visualizza.)
4. Fare clic sulla cella **Formato** per la mediana nell'elenco Visualizza e selezionare **nnnn** dall'elenco a discesa dei formati.
5. Nella cella Decimali, immettere 1.
6. Effettuare le stesse modifiche per la media nell'elenco Visualizza.
7. Clicca su **Applica a tutti** per applicare queste modifiche a tutte e tre le variabili di scala.
8. Clicca su **OK** nel builder da tavolo per creare la tabella.

	Mean	Median
Age of respondent	45.6	42.0
Highest year of school completed	13.3	13.0
Hours per day watching TV	2.9	2.0

Figura 39. Media e mediana visualizzate in tabella di variabili in scala

Conteggio, N valido e Valori mancanti

Spesso è utile visualizzare il numero di casi utilizzati per calcolare le statistiche di riepilogo, come la media, e si potrebbe supporre (non irragionevolmente) che la statistica di riepilogo *Conteggio* fornirebbe tali informazioni. Tuttavia, questo non vi darà una base di caso accurata se ci sono valori mancanti. Per ottenere una base di caso accurata, utilizzare *Valido N*.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).

2. Fare clic con il tasto destro del mouse su una qualsiasi delle variabili di scala in anteprima della tabella sul riquadro canvas e selezionare **Statistiche di riepilogo** dal menu a comparsa.
3. Nella finestra di dialogo Statistiche di riepilogo, selezionare **Conteggio** nell'elenco Statistiche e fare clic sulla freccia per aggiungendola alla lista Visualizza.
4. Quindi selezionare **N valido** nell'elenco Statistiche e fare clic sulla freccia per aggiungerlo all'elenco Visualizza.
5. Clicca su **Applica a tutti** per applicare queste modifiche a tutte e tre le variabili di scala.
6. Clicca su **OK** nel builder da tavolo per creare la tabella.

	Mean	Median	Count	Valid N
Age of respondent	45.6	42.0	2832	2828
Highest year of school completed	13.3	13.0	2832	2820
Hours per day watching TV	2.9	2.0	2832	2337

Figura 40. Conteggio versus Valid N

Per tutte e tre le variabili, *Conteggio* è lo stesso: 2.832. Non casualmente, questo è il numero totale di casi nel file dati. Poiché le variabili di scala non sono nidificate all'interno di eventuali variabili categoriali, *Conteggio* rappresenta semplicemente il numero totale di casi nel file dati.

Valido N, invece, è diverso per ogni variabile e differisce parecchio da *Conteggio* per *Ore al giorno guardando la TV*. Questo perché c'è un gran numero di **valori mancanti** per questa variabile -- ossia i casi senza valore registrati per questa variabile o valori definiti come rappresentare i dati mancanti (come un codice di 99 per rappresentare *Non Applicabile* per la gravidanza nei maschi).

7. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
8. Fare clic con il tasto destro del mouse su una qualsiasi delle variabili di scala in anteprima della tabella sul riquadro canvas e selezionare **Statistiche di riepilogo** dal menu a comparsa.
9. Nella finestra di dialogo Statistiche di riepilogo, selezionare **N valido** nell'elenco Visualizza e fare clic sul tasto freccia per spostarlo nuovamente nell'elenco Statistiche, rimuoverlo dall'elenco Visualizza.
10. Selezionare **Conteggio** nell'elenco Visualizza e fare clic sul tasto freccia per spostarlo nuovamente nell'elenco Statistiche, rimuoverlo dall'elenco Visualizza.
11. Selezionare **Manca** nell'elenco Statistiche e fare clic sul tasto freccia per aggiungerlo all'elenco Visualizza.
12. Clicca su **Applica a tutti** per applicare queste modifiche a tutte e tre le variabili di scala.
13. Clicca su **OK** nel builder da tavolo per creare la tabella.

	Mean	Median	Missing
Age of respondent	45.6	42.0	4
Highest year of school completed	13.3	13.0	12
Hours per day watching TV	2.9	2.0	495

Figura 41. Numero di valori mancanti visualizzati nella tabella delle statistiche di riepilogo della scala

La tabella visualizza ora il numero di valori mancanti per ogni variabile di scala. Questo rende abbastanza evidente che *Ore al giorno a guardare la TV* ha un gran numero di valori mancanti, mentre le altre due variabili ne hanno pochissime. Questo può essere un fattore da considerare prima di mettere una grande fede nei valori sommari per quella variabile.

Diversi Sommari per Variabili diverse

Oltre alla visualizzazione di più statistiche di riepilogo, è possibile visualizzare diverse statistiche di riepilogo per variabili di scala diverse in una tabella impilata. Ad esempio, la tabella precedente ha rivelato che solo una delle tre variabili ha un gran numero di valori mancanti; quindi si potrebbe voler mostrare il numero dei valori mancanti per solo quella variabile.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).

2. Clicca su *Age of respondent* nell'anteprima della tabella sul riquadro canvas, e poi Ctrl - clicca *Ultimo anno di scuola completato* per selezionare entrambe le variabili.
3. Fare clic con il tasto destro del mouse su una delle due variabili selezionate e selezionare **Statistiche di riepilogo** dal menu a comparsa.
4. Nella finestra di dialogo Statistiche di riepilogo, selezionare **Manca** nell'elenco Visualizza e fare clic sul tasto freccia per spostarlo nuovamente nell'elenco Statistiche, rimuoverlo dall'elenco Visualizza.
5. Clicca su **Applica alla selezione** per applicare la modifica solo alle due variabili selezionate.

I segnaposto nelle celle dati della tabella indicano che il numero dei valori mancanti verrà visualizzato solo per *Ore al giorno guardando la TV*.

6. Fare clic su **OK** per creare la tabella.

	Mean	Median	Missing
Age of respondent	45.6	42.0	
Highest year of school completed	13.3	13.0	
Hours per day watching TV	2.9	2.0	495

Figura 42. Tabella delle diverse statistiche di riepilogo per diverse variabili

Anche se questa tabella fornisce le informazioni che vogliamo, il layout potrebbe rendere difficile interpretare la tabella. Qualcuno che legge la tabella potrebbe pensare che le celle in bianco nella colonna *Manca* indicino zero valori mancanti per quelle variabili.

7. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
8. Nel gruppo di statistiche di riepilogo nel builder table, selezionare **Righe** dall'elenco a discesa Posizione.
9. Fare clic su **OK** per creare la tabella.

Age of respondent	Mean	45.6
	Median	42.0
Highest year of school completed	Mean	13.3
	Median	13.0
Hours per day watching TV	Mean	2.9
	Median	2.0
	Missing	495

Figura 43. Statistiche di riepilogo e variabili entrambe visualizzate nella dimensione della riga

Ora è chiaro che la tabella riporta il numero dei valori mancanti per una sola variabile.

Sommario Gruppo in Categorie

È possibile utilizzare variabili categoriali come variabili di raggruppamento per visualizzare riepiloghi variabili di scala all'interno di gruppi definiti dalle categorie della variabile categoriale.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Trascinare e rilasciare *Gender* dall'elenco delle variabili nell'area Colonne del riquadro canvas.

Se si fa clic con il tasto destro del mouse su *Gender* nell'anteprima della tabella sul riquadro canvas, si vedrà che **Statistiche di riepilogo** è disabilitata sul menu a comparsa. Questo perché in una tabella con variabili di scala, le variabili di scala sono sempre le variabili di origine statistiche.

3. Fare clic su **OK** per creare la tabella.

		Gender	
		Male	Female
Age of respondent	Mean	44.6	46.3
	Median	42.0	43.0
Highest year of school completed	Mean	13.4	13.2
	Median	13.0	13.0
Hours per day watching TV	Mean	2.8	2.9
	Median	2.0	2.0
	Missing	213	282

Figura 44. Riepiloghi su scala raggruppata utilizzando una variabile di colonna categoriale

Questo tavolo rende facile confrontare le medie (media e mediane) per maschi e femmine, e mostra chiaramente che tra loro non c'è molta differenza - che potrebbe non essere terribilmente interessante ma potrebbe essere un'informazione utile.

Variabili Raggruppamento Multipli

È possibile suddividere ulteriormente i gruppi per nidificazione e / o utilizzando variabili di raggruppamento categoriali di righe e colonne.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Trascinare e rilasciare *Ottieni notizie da internet* dall'elenco delle variabili al lato lontano dell'area Rows del riquadro canvas. Assicurati di posizionarlo in modo che tutte e tre le variabili di scala siano nidificate al suo interno, non solo una di esse.

Anche se ci possono essere momenti in cui vuoi qualcosa come il secondo esempio sopra, non è quello che vogliamo in questo caso.

3. Fare clic su **OK** per creare la tabella.

				Gender	
				Male	Female
Get news from internet	No	Age of respondent	Mean	47.0	48.8
			Median	45.0	46.0
		Highest year of school completed	Mean	13.4	13.1
			Median	13.0	12.0
		Hours per day watching TV	Mean	3.2	3.4
			Median	2.0	3.0
	Yes	Age of respondent	Missing	213	282
			Mean	38.7	41.1
			Median	35.0	38.0
		Highest year of school completed	Mean	13.2	13.3
			Median	13.0	13.0
		Hours per day watching TV	Mean	2.1	2.1
Median	2.0		2.0		
Missing	0		0		

Figura 45. Riepiloghi di scala raggruppati per variabili di righe e colonne categoriali

Nesting Categorical Variabili all'interno di Scale Variabili

Sebbene la tabella di cui sopra possa fornire le informazioni desiderate, potrebbe non fornirle nel formato più semplice da interpretare. Ad esempio, si può confrontare l'età media degli uomini che usano Internet per avere notizie e chi no - ma sarebbe più facile fare se i valori fossero prossimi l'uno all'altro piuttosto che separati. Scambiando le posizioni delle due variabili di riga e nidificando la variabile di raggruppamento categoriale all'interno delle tre variabili di scala potrebbe migliorare la tabella. Con le variabili di scala, il livello di nidificazione non ha alcun effetto sulla variabile di origine statistica. La variabile di scala è sempre la variabile di origine statistica indipendentemente dal livello di nidificazione.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Clicca su *Age of respondent* nella anteprima della tabella sul riquadro canvas, Ctrl - click *Ultime anno di scuola completato* e Ctrl - clicca *Ore al giorno guardando la TV* per selezionare tutte e tre le variabili di scala.

3. Trascinare e rilasciare le tre variabili di scala sul lato più lontano dell'area Rows, nidificando la variabile categoriale *Otteni notizie da internet* all'interno di ciascuna delle tre variabili di scala.
4. Fare clic su **OK** per creare la tabella.

				Gender	
				Male	Female
Age of respondent	Get news from internet	No	Mean	47.0	48.8
			Median	45.0	46.0
		Yes	Mean	38.7	41.1
			Median	35.0	38.0
Highest year of school completed	Get news from internet	No	Mean	13.4	13.1
			Median	13.0	12.0
		Yes	Mean	13.2	13.3
			Median	13.0	13.0
Hours per day watching TV	Get news from internet	No	Mean	3.2	3.4
			Median	2.0	3.0
			Missing	213	282
		Yes	Mean	2.1	2.1
			Median	2.0	2.0
			Missing	0	0

Figura 46. Variabile riga categoriale nidificata all'interno di variabili di scala impilate

La scelta dell'ordine di nidificazione dipende dalle relazioni o dai confronti che si desidera sottolineare in tavola. La modifica dell'ordine di nidificazione delle variabili di scala non modifica i valori delle statistiche di riepilogo; modifica solo le relative posizioni relative nella tabella.

Intervalli di confidenza

Gli intervalli di confidenza e gli errori standard sono disponibili per molte statistiche della tabella.

1. Dai menu, scegliere:

Analisi > Tavole > Tabelle personalizzate ...

2. Nel builder table, spostare *Highest degree* nell'area di riga del riquadro canvas.
3. Fare clic su **Statistiche di riepilogo ...**
4. Nell'elenco **Statistiche** nella finestra di dialogo **Statistiche di riepilogo**, fare clic sull'icona accanto a **Conteggio** per espandere l'elenco delle statistiche relative al conteggio.
5. Muovi le seguenti statistiche all'area **Visualizza**: Colonna N%, Bassa CL per il conteggio delle colonne%, Alto CL per il conteggio delle colonne% e errore standard del conteggio delle colonne.
6. Nel gruppo **intervalli di confidenza**, per **Livello (%)**, inserire 99.
7. Fare clic su **Applica alla selezione**, quindi fare clic su **Chiudi**.
8. Fare clic su **OK** per creare la tabella.

		Count	Column N %	99.0% Lower CL for Column N %	99.0% Upper CL for Column N %	Standard Error of Column N %
Highest degree	LT High school	430	15.2%	13.6%	17.0%	0.7%
	High school	1500	53.2%	50.7%	55.6%	0.9%
	Junior college	209	7.4%	6.2%	8.7%	0.5%
	Bachelor	478	16.9%	15.2%	18.8%	0.7%
	Graduate	205	7.3%	6.1%	8.6%	0.5%

Figura 47. Tabella dei conteggi, percentuali di colonna e intervalli di confidenza

9. Ricordare la finestra di dialogo **Tabelle personalizzate** e fare clic su **Statistiche di riepilogo ...**
10. Per il limite di confidenza inferiore, modificare l'etichetta in "Bassa di fiducia (& [livello di confidenza])". La stringa "& [Confidence Level]" inserisce il valore del livello di confidenza specificato in quella sede in etichetta.

11. Fare clic su **Applica alla selezione**, quindi fare clic su **Chiudi**.
12. Fare clic su **OK** per creare la tabella.

		Count	Column N %	Lower Confidence Limit (99.0%)	99.0% Upper CL for Column N %	Standard Error of Column N %
Highest degree	LT High school	430	15.2%	13.6%	17.0%	0.7%
	High school	1500	53.2%	50.7%	55.6%	0.9%
	Junior college	209	7.4%	6.2%	8.7%	0.5%
	Bachelor	478	16.9%	15.2%	18.8%	0.7%
	Graduate	205	7.3%	6.1%	8.6%	0.5%

Figura 48. Tabella con etichetta intervallo di confidenza modificato

Statistiche del test

Statistiche del test

Sono disponibili tre diversi test di significati per studiare la relazione tra variabili di righe e colonne. Questo capitolo discute l'output di ciascuno di questi test, con particolare attenzione agli effetti di nidificazione e impilamento. Per ulteriori informazioni, consultare l'argomento [“Stack, Nesting e Layers con Variabili categoriali”](#) a pagina 24.

File di dati di esempio

Gli esempi in questo capitolo utilizzano il file di dati *survey_sample.sav*. Per ulteriori informazioni, consultare l'argomento [“File di esempio”](#) a pagina 76.

Test di indipendenza (Chi-quadrato)

Il test di indipendenza del chi - quadrato viene utilizzato per determinare se esiste una relazione tra due variabili categoriali. Ad esempio, si potrebbe voler stabilire se *Stato forza lavoro* è correlato a *stato civile*.

1. Dai menu, scegliere:

Analisi > Tavole > Tabelle personalizzate ...

2. Nel builder table, trascinare e rilasciare *Stato forza lavoro* dall'elenco di variabili nell'area Rows del riquadro canvas.
3. Trascinare e rilasciare *stato civile* dall'elenco delle variabili nell'area Colonne.
4. Selezionare **Righe** come posizione per le statistiche di riepilogo.
5. Selezionare *Stato forza lavoro* e fare clic su **Statistiche di riepilogo** nel gruppo Definisci.
6. Selezionare **Colonna N%** nell'elenco Statistiche e aggiungelo alla lista Visualizza.
7. Fare clic su **Applica alla selezione**.
8. Nella finestra di dialogo Tabelle personalizzate, fare clic sulla scheda **Statistiche di test**.
9. Selezionare **Test di indipendenza (Chi-square)**.
10. Clicca su **OK** per creare la tabella e ottenere il test chi quadrato.

			Marital status				
			Married	Widowed	Divorced	Separated	Never married
Labor force status	Working full time	Count	778	44	295	58	392
		Column %	57.8%	15.5%	66.1%	62.4%	59.1%
	Working part-time	Count	138	20	35	9	102
		Column %	10.3%	7.1%	7.8%	9.7%	15.4%
	Temporarily not working	Count	23	2	9	1	11
		Column %	1.7%	.7%	2.0%	1.1%	1.7%
	Unemployed, laid off	Count	13	3	10	0	32
		Column %	1.0%	1.1%	2.2%	.0%	4.8%
	Retired	Count	168	150	53	6	17
		Column %	12.5%	53.0%	11.9%	6.5%	2.6%
	School	Count	9	1	7	2	60
		Column %	.7%	.4%	1.6%	2.2%	9.0%
	Keeping house	Count	200	55	25	13	35
		Column %	14.9%	19.4%	5.6%	14.0%	5.3%
	Other	Count	16	8	12	4	14
		Column %	1.2%	2.8%	2.7%	4.3%	2.1%

Figura 49. Stato forza lavoro per stato civile

Questa tabella è una crosstabulazione di *Stato forza lavoro* da *stato Marital*, con conteggi e proporzioni colonne mostrate come statistiche di riepilogo. Le proporzioni della colonna sono calcolate in modo che somma al 100% in basso ciascuna colonna. Se queste due variabili non sono correlate, allora in ogni riga le proporzioni dovrebbero essere simili tra le colonne. Sembrano esserci differenze nelle proporzioni, ma si può verificare il test chi quadrato per essere sicuri.

		Marital status
Labor force status	Chi-square	729.242
	df	28
	Sig.	.000*

*. The Chi-square statistic is significant at the 0.05 level.

Figura 50. Il test di chi - quadrato di Pearson

Il test di indipendenza ipotizza che *Labor force status* e *Stato Marital* siano indipendenti - ovvero che le proporzioni della colonna siano le stesse su colonne, e le eventuali discrepanze osservate siano dovute a variazioni di possibilità. La statistica chi quadrato misura la discrepanza complessiva tra la conta delle cellule osservate e i conteggi che ci si aspetterebbe se le proporzioni della colonna fossero le stesse attraverso le colonne. Una statistica chi - quadrato più grande indica una maggiore discrepanza tra la conta delle cellule osservate e attese - evidenza maggiore che le proporzioni della colonna non sono uguali, che l'ipotesi di indipendenza è errata e, quindi, che *Stato forza lavoro* e *stato civile* sono correlati.

La statistica chi - quadrato calcolata ha un valore di 729.242. Per stabilire se si tratta di prove sufficienti per respingere l'ipotesi di indipendenza, viene calcolato il valore di significatività della statistica. Il valore di significatività è la probabilità che una variata casuale estratta da una distribuzione chi - quadrato con 28 gradi di libertà sia maggiore di 729.242. Poiché questo valore è inferiore al livello alfa specificato nella scheda Statistiche test, è possibile rifiutare l'ipotesi di indipendenza al livello 0.05 . Così, *Stato forza lavoro* e *stato civile* sono in realtà correlati.

Effetti di Nesting e Stacking su Testi di Indipendenza

La regola per i test di indipendenza è la seguente: viene eseguita una prova separata per ogni sottotetto più innervi. Per vedere come la nidificazione influenzi i test, considerare l'esempio precedente, ma con *stato Marital* nidificato all'interno dei livelli di *Gender*.

1. Aprire di nuovo il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Trascinare e rilasciare *Gender* dall'elenco delle variabili nell'area Colonne del riquadro canvas sopra *stato civile*.
3. Fare clic su **OK** per creare la tabella.

		Gender	
		Male	Female
		Marital status	Marital status
Labor force status	Chi-square	246.637	542.589
	df	28	28
	Sig.	.000*.1,2	.000*.1,2

*. The Chi-square statistic is significant at the 0.05 level.

1. More than 20% of cells in this sub-table have expected cell counts less than 5.

2. The minimum expected cell count in this sub-table is less than one.

Figura 51. Il test di chi - quadrato di Pearson

Con *Stato Marital* nidificato all'interno dei livelli di *Gender* vengono eseguiti due test - uno per ogni livello di *Gender*. Il valore di significatività per ogni prova indica che è possibile rifiutare l'ipotesi di indipendenza tra *status matrimoniale* e *Stato forza lavoro* sia per i maschi che per le femmine. Tuttavia, la tabella rileva che più del 20% delle celle di ciascuna tabella hanno conteggi inferiori a 5, e il conteggio delle celle minime previsto è inferiore a 1. Queste note indicano che le ipotesi del test di chi - quadrato non possono essere soddisfatte da queste tabelle, quindi i risultati dei test sono sospetti.

Nota: le note a piè di pagina possono essere tagliate fuori vista dai confini cellulari. È possibile renderli visibili modificando l'allineamento di queste celle nella finestra di dialogo Proprietà Cell.

Per vedere come gli stack influiscono sui test:

4. Aprire di nuovo il builder table (Analizza menu, tabelle, Tabelle personalizzate).
5. Trascinare e rilasciare *Altezza più alta* dall'elenco delle variabili nell'area Rows sotto *Stato forza lavoro*.
6. Fare clic su **OK** per creare la tabella.

		Gender	
		Male	Female
		Marital status	Marital status
Labor force status	Chi-square	246.637	542.589
	df	28	28
	Sig.	.000*.1,2	.000*.1,2
Highest degree	Chi-square	43.844	105.506
	df	16	16
	Sig.	.000*	.000*

*. The Chi-square statistic is significant at the 0.05 level.

1. More than 20% of cells in this sub-table have expected cell counts less than 5.

2. The minimum expected cell count in this sub-table is less than one.

Figura 52. Il test di chi - quadrato di Pearson

Con *Ultimo grado* impilati con *Stato forza lavoro*, vengono eseguiti quattro test - un test di indipendenza di *Stato Marital* e *Stato forza lavoro* e una prova di *stato civile* e *Grado più alto* per ogni livello di *Gender*. I risultati del test per *stato Marital* e *Stato forza lavoro* sono gli stessi di prima. I risultati del test per *Stato Marital* e *Highest degree* indicano queste variabili non sono indipendenti.

Confronto Colonna Significa

La colonna significa i test utilizzati per determinare se vi sia una relazione tra una variabile categoriale nelle Colonne e una variabile continua nelle righe. Inoltre, è possibile utilizzare i risultati del test per determinare l'ordinamento relativo di categorie della variabile categoriale in termini di valore medio della variabile continua. Ad esempio, si potrebbe voler stabilire se *Ore al giorno guardando la TV* è correlato a *Get news dai giornali*.

1. Dai menu, scegliere:

Analisi > Tavole > Tabelle personalizzate ...

2. Clicca su **Reset** per ripristinare le impostazioni predefinite a tutte le schede.

3. Nel builder table, trascinare e rilasciare *Ore al giorno guardando la TV* dall'elenco delle variabili nell'area Rows del riquadro canvas.

4. Trascinare e rilasciare *Ottieni notizie dai giornali* dall'elenco delle variabili nell'area Colonne.

5. Selezionare *Ore al giorno guardando la TV* e fare clic su **Statistiche di riepilogo** nel gruppo Definisci.

6. Selezionare **nnnn** come formato.

7. Selezionare **2** come numero di decimali da visualizzare. Notare che ciò fa sì che il formato ora legga **nnnn.nn**.

8. Fare clic su **Applica alla selezione**.

9. Nella finestra di dialogo Tabelle personalizzate, fare clic sulla scheda **Statistiche di test**.

10. Selezionare **Confronta mezzi di colonna (t - test)**.

11. Fare clic su **OK** per creare la tabella e ottenere la colonna significa test.

	Get news from newspapers	
	No	Yes
	Mean	Mean
Hours per day watching TV	2.92	2.74

Figura 53. Ricevi notizie dai giornali di Ore al giorno guardando la TV

In questa tabella è riportata la media *Ore al giorno guardando la TV* per le persone che fanno e non ottengono la loro notizia dai giornali. La differenza osservata in questi mezzi suggerisce che le persone che non ricevono le loro notizie dai giornali trascorrono circa 0.18 ore in più a guardare la TV rispetto alle persone che ricevono le loro notizie dai giornali. Per verificare se questa differenza è dovuta alla variazione di possibilità, controllare la colonna significa test.

	Get news from newspapers	
	No	Yes
	(A)	(B)
Hours per day watching TV		

Figura 54. Confronti delle medie di colonna

La tabella dei mezzi di prova assegna una chiave di lettera a ciascuna categoria della variabile di colonna. Per *Ottieni notizie dai giornali*, alla categoria *No* viene assegnata la lettera A e *Sì* viene assegnata la lettera B. Per ogni coppia di colonne, i mezzi di colonna vengono confrontati utilizzando un test *t*. Poiché ci sono solo due colonne, viene eseguito un solo test. Per ogni coppia significativa, la chiave della categoria con la media più piccola viene inserita sotto la categoria con media maggiore. Dal momento che non sono riportate chiavi nelle celle della tabella, ciò significa che la colonna significa non è statisticamente diversa.

Rilevanza Risultati in Notazione stile APA

Se non si desidera che il significato risulti in una tabella separata, è possibile scegliere di visualizzarle nella tabella principale. I risultati di significati sono identificati utilizzando una notazione in stile APA con lettere di sottoscrittura. Completare i passi precedenti per confrontare i mezzi di colonna, ma effettuare la seguente modifica nella scheda Statistiche di test:

1. Nell'area Identificazione Differenze Significative, selezionare **Nella tabella principale utilizzando gli script secondari di stile APA**.

2. Fare clic su **OK** per creare la tabella e ottenere la colonna significa test tramite notazione in stile APA.

	Get news from newspapers	
	No	Yes
	Mean	Mean
Hours per day watching TV	2.92 _a	2.74 _a

Figura 55. Confronti di colonna significa utilizzare notazione in stile APA

La tabella dei mezzi di prova assegna una lettera di sottoscrizione alle categorie della variabile di colonna. Per ogni coppia di colonne, i mezzi di colonna vengono confrontati utilizzando un test t . Se una coppia di valori è significativamente differente, i valori hanno *diverse* lettere di sottoscrizione assegnate. Poiché ci sono solo due colonne, viene eseguito un solo test. Poiché la colonna significa in questo esempio condividere la stessa lettera di sottoscrizione, la colonna significa non è statisticamente diversa.

Effetti di Nesting e Stacking su Colonna Medie Prove

La regola per la colonna significa i test è la seguente: viene eseguita una serie separata di test pairwise per ogni sottotetto più innerabile. Per vedere come la nidificazione influenzi i test, considerare l'esempio precedente, ma con *Ore al giorno guardando la TV* nidificata all'interno dei livelli di *Stato forza lavoro*.

1. Aprire di nuovo il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Trascinare e rilasciare *Stato forza lavoro* dall'elenco delle variabili nell'area Rows del riquadro canvas.
3. Fare clic su **OK** per creare la tabella.

			Get news from newspapers	
			No	Yes
			(A)	(B)
Labor force status	Working full time	Hours per day watching TV	B	
	Working part-time	Hours per day watching TV		
	Temporarily not working	Hours per day watching TV		
	Unemployed, laid off	Hours per day watching TV		
	Retired	Hours per day watching TV		
	School	Hours per day watching TV		
	Keeping house	Hours per day watching TV		
	Other	Hours per day watching TV		

Figura 56. Confronti delle medie di colonna

Con *Ore al giorno a guardare la TV* nidificata all'interno dei livelli di *Stato forza lavoro*, vengono eseguite sette serie di colonne medie: una per ogni livello di *Stato forza lavoro*. Le stesse chiavi di lettera vengono assegnate alle categorie di *Richiamo notizie dai giornali*. Per gli intervistati *lavorare a tempo pieno*, la chiave B compare nella colonna A. Questo significa che per i dipendenti a tempo pieno, il valore medio di *Ore al giorno a guardare la TV* è più basso per le persone che ottengono la loro notizia dai giornali. Non compaiono altre chiavi nelle colonne, quindi si può concludere che non ci siano altre differenze statisticamente significative nella colonna mezzi.

aggiustamenti Bonferroni. Quando vengono eseguite più prove, la regolazione Bonferroni viene applicata alla colonna significa test per garantire che il livello alfa (o falso tasso positivo) specificato sulla scheda Statistiche del test si applichi a ciascun *set* di test. Così, in questa tabella, non sono stati applicati aggiustamenti Bonferroni perché nonostante siano state eseguite sette serie di test, all'interno di ciascuna serie viene confrontato un solo paio di colonne.

Per vedere come gli stack influiscono sui test:

4. Aprire di nuovo il builder table (Analizza menu, tabelle, Tabelle personalizzate).
5. Trascinare e rilasciare *Ottieni notizie da internet* dall'elenco delle variabili nell'area Colonne a sinistra di *Ottieni notizie dai giornali*.
6. Fare clic su **OK** per creare la tabella.

			Get news from internet		Get news from newspapers	
			No	Yes	No	Yes
			(A)	(B)	(A)	(B)
Labor force status	Working full time	Hours per day watching TV	B		B	
	Working part-time	Hours per day watching TV	B			
	Temporarily not working	Hours per day watching TV				
	Unemployed, laid off	Hours per day watching TV	B			
	Retired	Hours per day watching TV	B			
	School	Hours per day watching TV	B			
	Keeping house	Hours per day watching TV	B			
	Other	Hours per day watching TV	B			

Figura 57. Confronti delle medie di colonna

Con *Ricevi notizie da internet* impilate con *Ricevi notizie dai giornali*, vengono eseguite 14 serie di test di colonna: uno per ogni livello di *Stato forza lavoro* per *Ricevi notizie da internet* e *Ricevi notizie dai giornali*. Anche in questo caso non vengono applicati aggiustamenti Bonferroni perché all'interno di ogni set viene comparata una sola coppia di colonne. I test per *Get news dai giornali* sono gli stessi di prima. Per *Ottieni notizie da internet*, alla categoria *No* viene assegnata la lettera A e *Sì* viene assegnata la lettera B. La chiave B è riportata nella colonna A per ogni serie di colonne significa test tranne che per quegli intervistati temporaneamente non funzionanti. Questo significa che il valore medio di *Ore al giorno a guardare la TV* è più basso per le persone che ottengono la loro notizia da Internet che per le persone che non ottengono la loro notizia dai giornali. Non sono riportati chiavi per il set *Temporaneamente non funzionanti*; quindi, i mezzi di colonna non sono statisticamente diversi per questi intervistati.

Confronto Proporzioni Colonne

Le prove delle proporzioni della colonna sono utilizzate per determinare l'ordinamento relativo di categorie della variabile categoriale delle Colonne in termini di proporzioni di categoria della variabile categoriale Rows. Ad esempio, dopo aver utilizzato un test di chi - quadrato per scoprire che *Stato forza lavoro* e *stato Marital* non sono indipendenti, è possibile scoprire quali righe e colonne sono responsabili di questa relazione.

1. Dai menu, scegliere:

Analisi > Tavole > Tabelle personalizzate ...

2. Clicca su **Reset** per ripristinare le impostazioni predefinite a tutte le schede.
3. Nel builder table, trascinare e rilasciare *Stato forza lavoro* dall'elenco di variabili nell'area Rows del riquadro canvas.
4. Trascinare e rilasciare *stato civile* dall'elenco delle variabili nell'area Colonne.
5. Selezionare *Stato forza lavoro* e fare clic su **Statistiche di riepilogo** nel gruppo Definisci.
6. Selezionare **Colonna N%** nell'elenco Statistiche e aggiungelo alla lista Visualizza.
7. Deselezionare **Conteggio** dall'elenco Visualizza.
8. Fare clic su **Applica alla selezione**.

9. Nella finestra di dialogo Tabelle personalizzate, fare clic sulla scheda **Statistiche di test**.
10. Selezionare **Confronta proporzioni colonne (z - test)**.
11. Clicca su **OK** per creare la tabella e ottenere i test delle proporzioni della colonna.

		Marital status				
		Married	Widowed	Divorced	Separated	Never married
		Column %	Column %	Column %	Column %	Column %
Labor force status	Working full time	57.8%	15.5%	66.1%	62.4%	59.1%
	Working part-time	10.3%	7.1%	7.8%	9.7%	15.4%
	Temporarily not working	1.7%	.7%	2.0%	1.1%	1.7%
	Unemployed, laid off	1.0%	1.1%	2.2%	.0%	4.8%
	Retired	12.5%	53.0%	11.9%	6.5%	2.6%
	School	.7%	.4%	1.6%	2.2%	9.0%
	Keeping house	14.9%	19.4%	5.6%	14.0%	5.3%
	Other	1.2%	2.8%	2.7%	4.3%	2.1%

Figura 58. Stato forza lavoro per stato civile

Questa tabella è una crosstabulazione di *Stato forza lavoro* da *stato Marital*, con proporzioni colonne mostrate come statistica di riepilogo.

		Marital status				
		Married	Widowed	Divorced	Separated	Never married
		(A)	(B)	(C)	(D)	(E)
Labor force status	Working full time	B		A B	B	B
	Working part-time					A B C
	Temporarily not working					
	Unemployed, laid off					A B
	Retired	E	A C D E	E		
	School					A B C
	Keeping house	C E	C E		C E	
	Other					

Figura 59. Confronti delle proporzioni di colonna

La tabella delle proporzioni della colonna assegna una chiave di lettera a ciascuna categoria delle variabili della colonna. Per *stato civile*, alla categoria *Married* viene assegnata la lettera A, *Widowed* viene assegnata la lettera B, e così via, attraverso la categoria *Mai sposati*, a cui viene assegnata la lettera E. Per ogni coppia di colonne, le proporzioni della colonna vengono confrontati utilizzando un test z. Vengono eseguite sette serie di prove di proporzioni della colonna, una per ogni livello di *Stato forza lavoro*. Poiché ci sono cinque livelli di *stato Marital*, $(5 * 4) / 2 = 10$ coppie di colonne sono confrontati in ciascuna serie di test, e le regolazioni Bonferroni sono utilizzate per regolare i valori di significatività. Per ogni coppia significativa, la chiave della categoria più piccola viene inserita sotto la categoria con la proporzione maggiore.

Per la serie di test associati a *Lavorare a tempo pieno*, compare la chiave B in ognuna delle altre colonne. Inoltre, la chiave A compare nella colonna C. Non sono riportate altre chiavi in altre colonne. Si può concludere che la percentuale di persone divorziate che lavorano a tempo pieno è superiore alla proporzione di persone sposate che lavorano a tempo pieno, che a sua volta è maggiore della proporzione di vedovi che lavorano a tempo pieno. Le proporzioni di persone separate o mai sposate e lavorate a tempo pieno non possono essere differenziate da persone divorziate o sposate e che lavorano a tempo pieno, ma queste proporzioni sono superiori alla proporzione di vedovi che lavorano a tempo pieno.

Per i test associati a *Tempo di lavoro o Scuola*, le chiavi A, B e C appaiono nella colonna E. Non sono riportate altre chiavi in altre colonne. Così, le proporzioni di persone che non si sono mai sposate e sono a scuola o che lavorano part time sono superiori alle proporzioni di persone sposate, vedove o divorziate che si trovano a scuola o lavorano part time.

Per i test associati a *Temporalmente non funzionanti* o con lo stato di lavoro *Altro*, non vengono riportate altre chiavi in nessuna colonna. Così, non vi è alcuna differenza discernibile nelle proporzioni di persone sposate, vedove, divorziate, separate o neverne che temporaneamente non lavorano o si trovano in una situazione di occupazione altrimenti non categorizzata.

I test associati a *Pensionati* mostrano che la percentuale di vedovi che si ritirano è superiore alle proporzioni di tutte le altre categorie matrimoniali ritirate. Inoltre, le proporzioni di persone sposate o divorziate che si ritirano sono superiori alla percentuale di persone neosposate che si ritirano.

Ci sono proporzioni maggiori di persone sposate, vedove o separate e che tengono casa che proporzioni di persone divorziate o mai sposate e che tengono casa.

La percentuale di persone che non si sono mai sposate e sono *Non occupati, licenziati* è superiore alle proporzioni di persone sposate o vedove e disoccupate. Inoltre, notare che la colonna *Separata* è contrassegnata con un ".", che indica che la percentuale osservata di persone separate nella riga *Unoccupati, licenziati* è di 0 o 1 e quindi non possono essere effettuati confronti utilizzando quella colonna per gli intervistati disoccupati.

Fusione Significato Risultati nella Tabella Principale

Se non si desidera che il significato risulti in una tabella separata, è possibile scegliere di visualizzarle nella tabella principale. Completa i passi precedenti per confrontare le proporzioni della colonna, ma effettuare la seguente modifica nella scheda Statistiche di test:

1. Nell'area Identificazione Differenze Significative, selezionare **Nella tabella principale**.
2. Fare clic su **OK** per creare la tabella.

		Marital status				
		Married	Widowed	Divorced	Separated	Never married
		(a)	(b)	(c)	(d)	(e)
		Column N %	Column N %	Column N %	Column N %	Column N %
Labor force status	Working full time	57.8% b	15.5%	66.1% a b	62.4% b	59.1% b
	Working part-time	10.3%	7.1%	7.8%	9.7%	15.4% a b c
	Temporarily not working	1.7%	0.7%	2.0%	1.1%	1.7%
	Unemployed, laid off	1.0%	1.1%	2.2%	0.0% ¹	4.8% a b
	Retired	12.5% e	53.0% a c d e	11.9% e	6.5%	2.6%
	School	0.7%	0.4%	1.6%	2.2%	9.0% a b c
	Keeping house	14.9% c e	19.4% c e	5.6%	14.0% c e	5.3%
	Other	1.2%	2.8%	2.7%	4.3%	2.1%

Figura 60. Risultati di significati uniti nella tabella principale

3. Fare doppio clic sulla tabella nella finestra **Viewer** per attivarla.
4. Passare sopra una delle celle delle etichette delle colonne che contengono le chiavi della lettera. Per esempio, passare sopra la cella con l'etichetta "(a)".

		Marital status				
		Married	Widowed	Divorced	Separated	Never married
		(a)	(b)	(c)	(d)	(e)
		Column N %	Column N %	Column N %	Column N %	Column N %
Labor force status	Working full time	57.8% b	15.5%	66.1% a b	62.4% b	59.1% b
	Working part-time	10.3%	7.1%	7.8%	9.7%	15.4% a b c
	Temporarily not working	1.7%	0.7%	2.0%	1.1%	1.7%
	Unemployed, laid off	1.0%	1.1%	2.2%	0.0% ¹	4.8% a b
	Retired	12.5% e	53.0% a c d e	11.9% e	6.5%	2.6%
	School	0.7%	0.4%	1.6%	2.2%	9.0% a b c
	Keeping house	14.9% c e	19.4% c e	5.6%	14.0% c e	5.3%
	Other	1.2%	2.8%	2.7%	4.3%	2.1%

Figura 61. Evidenziazione di tutte le celle che contengono la chiave di lettera selezionata

Tutte le celle che contengono quella chiave sono evidenziate nella tabella.

5. Clicca con il tasto destro del mouse sulla cella etichetta della colonna. Dal menu contestuale scegliere **Seleziona > Selezionare tutte le celle con questa chiave di significato.**

		Marital status				
		Married	Widowed	Divorced	Separated	Never married
		(a)	(b)	(c)	(d)	(e)
		Column N %	Column N %	Column N %	Column N %	Column N %
Labor force status	Working full time	57.8% b	15.5%	66.1% a b	62.4% b	59.1% b
	Working part-time	10.3%	7.1%	7.8%	9.7%	15.4% a b c
	Temporarily not working	1.7%	0.7%	2.0%	1.1%	1.7%
	Unemployed, laid off	1.0%	1.1%	2.2%	0.0% ¹	4.8% a b
	Retired	12.5% e	53.0% a c d e	11.9% e	6.5%	2.6%
	School	0.7%	0.4%	1.6%	2.2%	9.0% a b c
	Keeping house	14.9% c e	19.4% c e	5.6%	14.0% c e	5.3%
	Other	1.2%	2.8%	2.7%	4.3%	2.1%

Figura 62. Selezione di tutte le celle che contengono la chiave di lettera selezionata

Tutte le celle che contengono quella chiave sono selezionate nella tabella.

Effetti di Nesting e Stacking su Colonna Proporzioni Test

La regola per i test delle proporzioni della colonna è la seguente: viene eseguita una serie separata di test pairwise per ogni sottotetto più innerabile. Per vedere come la nidificazione influenzi i test, considerare l'esempio precedente, ma con *Labor force status* nidificato all'interno dei livelli di *Gender*.

1. Aprire di nuovo il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Trascinare e rilasciare *Gender* dall'elenco delle variabili nell'area Rows del riquadro canvas.
3. Fare clic su **OK** per creare la tabella.

				Marital status				
				Married	Widowed	Divorced	Separated	Never married
				(A)	(B)	(C)	(D)	(E)
Gender	Male	Labor force status	Working full time	B		B	B	B
			Working part-time					A
			Temporarily not working		.			
			Unemployed, laid off		.			A
			Retired	E	A C D E	E		
			School		.			A C
			Keeping house					
	Female	Labor force status	Other				A	
			Working full time	B		A B	B	B
			Working part-time	B				B
			Temporarily not working					
			Unemployed, laid off					A
			Retired	E	A C D E	E		
			School					A B C
			Keeping house	C E	C E		C	
			Other					

Figura 63. Confronti delle proporzioni di colonna

Con *stato di forza Labor* nidificato all'interno dei livelli di *Gender*, vengono eseguite 14 serie di test delle proporzioni della colonna -- una per ogni livello di *Stato forza lavoro* per ogni livello di *Gender*. Le stesse chiavi di lettera sono assegnate alle categorie di *stato civile*.

Ci sono un paio di cose da notare sui risultati della tavola:

- Con più test ci sono più colonne con quota zero colonne. Sono più comuni tra gli intervistati separati e i maschi vedenti.
- Le differenze di colonna viste in precedenza tra gli intervistati *tenere casa* sembrano essere interamente dovute alle femmine.

Per vedere come gli stack influiscono sui test:

4. Aprire di nuovo il builder table (Analizza menu, tabelle, Tabelle personalizzate).
5. Trascinare e rilasciare *Grado più alto* dall'elenco di variabili nell'area Rows sotto *Gender*.
6. Fare clic su **OK** per creare la tabella.

				Marital status					
				Married	Widowed	Divorced	Separated	Never married	
				(A)	(B)	(C)	(D)	(E)	
Gender	Male	Labor force status	Working full time	B		B	B	B	
			Working part-time					A	
			Temporarily not working						
			Unemployed, laid off					A	
			Retired	E	A C D E	E			
			School					A C	
			Keeping house						
	Female	Labor force status	Other				A		
			Working full time	B		A B	B	B	
			Working part-time	B				B	
			Temporarily not working						
			Unemployed, laid off					A	
			Retired	E	A C D E	E			
			School					A B C	
Highest degree	Keeping house			C E	C E		C		
				Other					
	LT High school							A C E	
	High school								
	Junior college			B		B		B	
	Bachelor			B				B	
	Graduate			B					

Figura 64. Confronti delle proporzioni di colonna

Con *Highest degree* sovrapposto a *Gender*, vengono eseguite 19 serie di test di colonna - i 14 precedentemente discussi più uno per ogni livello di *Highest degree*. Le stesse chiavi di lettera sono assegnate alle categorie di *stato civile*.

Ci sono alcune cose da notare sui risultati della tabella:

- I risultati di prova per le 14 serie di prove precedentemente eseguite sono gli stessi.
- Le persone che hanno meno di un diploma di scuola superiore sono più comuni tra i vedenti che tra gli intervistati sposati, divorziati o neo - sposati.
- Le persone con qualche istruzione scolastica post - alta tendono ad essere più comuni tra quelle persone che sono sposate, divorziate e mai sposate che tra i vedovi.

Una nota sui pesi e gli insiemi a risposta multipla

I pesi dei casi sono sempre basati su conteggi, non risposte, anche quando una delle variabili è una variabile di risposta multipla.

Insiemi a risposta multipla

Le tabelle personalizzate e il Builder di grafico supportano un tipo speciale di "variabile" detta **insieme a risposta multipla**. Gli insiemi a risposta multipla non possono essere considerati variabili nell'accezione propria del termine, poiché non è possibile visualizzarli nell'Editor dei dati e non vengono riconosciuti dalle altre procedure. Gli insiemi a risposta multipla utilizzano più variabili per registrare le risposte alle domande in cui il rispondente può fornire più di una risposta. Tali insiemi vengono considerati come variabili categoriali e nella maggior parte dei casi consentono di eseguire le stesse operazioni previste per questo tipo di variabili.

Gli insiemi a risposta multipla vengono creati utilizzando più variabili del file di dati e rappresentano un tipo di insieme speciale all'interno di un file di dati. È possibile definire e salvare più insiemi a risposta multipla in un file di dati in formato IBM SPSS Statistica, ma non è possibile importare o esportare tali insiemi da o in altri formati di file. È possibile copiare gli insiemi a risposta multipla da altri file di dati di IBM SPSS Statistica tramite l'opzione Copia proprietà dati accessibile tramite il menu Dati nella finestra Editor dei dati.

File di dati di esempio

Gli esempi in questo capitolo utilizzano il file di dati *survey_sample.sav*. Per ulteriori informazioni, consultare l'argomento "File di esempio" a pagina 76.

Tutti gli esempi forniti qui visualizzano etichette variabili in finestre di dialogo, ordinate in ordine alfabetico. Le proprietà di visualizzazione di elenco variabili sono specificate nella scheda Generale nella finestra di dialogo Opzioni (Modifica menu, Opzioni).

Conteggi, Risposte, Percentuali e Totali

Tutte le statistiche di riepilogo disponibili per le variabili categoriali sono disponibili anche per più set di risposta. Alcune statistiche aggiuntive sono disponibili anche per più set di risposta.

1. Dai menu, scegliere:

Analisi > Tavole > Tabelle personalizzate ...

2. Trascinare e rilasciare le *Fonti News* (questa è l'etichetta descrittiva per la serie di risposta multipla *\$mltnews*) dall'elenco delle variabili nell'area Rows del riquadro canvas.

L'icona accanto alla "variabile" nell'elenco variabili lo identifica come una serie di dicotomia multipla.



Figura 65. Icona serie di dicotomia multipla

Per una serie di dicotomia multipla, ogni "categoria" è, infatti, una variabile separata e le etichette di categoria sono le etichette variabili (o i nomi variabili per le variabili senza etichette variabili definite). In questo esempio, i conteggi che verranno visualizzati rappresentano il numero di casi con una risposta Sì per ciascuna variabile nel set.

3. Fare clic con il tasto destro del mouse su *Origini notizie* nell'anteprima della tabella sul riquadro canvas e selezionare **Categorie e Totali** dal menu a comparsa.
4. Selezionare (cliccare) **Totale** nella finestra di dialogo Categorie e Totali, quindi fare clic su **Applica**.
5. Fare clic con il tasto destro del mouse su *News* e selezionare **Statistiche di riepilogo** dal menu a comparsa.
6. Nella finestra di dialogo Statistiche di riepilogo, selezionare **Colonna N%** nell'elenco Statistiche e fare clic sulla freccia per aggiungendola alla lista Visualizza.
7. Fare clic su **Applica alla selezione**, quindi fare clic su **OK** per creare la tabella.

		Count	Column N %
News sources	Get news from internet	867	41.7%
	Get news from radio	551	26.5%
	Get news from television	1077	51.8%
	Get news from news magazines	294	14.1%
	Get news from newspapers	805	38.7%
	Total	2081	100.0%

Figura 66. Contatori di dicotomia multipla e percentuali di colonna

Totali che non Aggiungi Up

Se si guarda ai numeri della tabella, si può notare che c'è una discrepanza abbastanza ampia tra i "totali" e i valori presumibilmente totalizzati - nello specifico, i totali sembrano essere molto più bassi di quanto

dovrebbero essere. Questo perché il conteggio per ogni "categoria" nella tabella è il numero di casi con un valore di 1 (una risposta Sì) per quella variabile e il numero totale di risposte Sì per tutte e cinque le variabili nel set di dicotomia multipla potrebbe facilmente superare il numero totale di casi nel file dati.

Il totale "conta", tuttavia, è il numero totale di casi con risposta Sì per almeno una variabile nel set, che non può mai superare il numero totale di casi nel file dati. In questo esempio, il conteggio totale di 2.081 è inferiore di quasi il 800 rispetto al totale dei casi nel file dei dati. Se nessuna di queste variabili ha valori mancanti, questo significa che quasi 800 intervistati intervistati hanno indicato di non ricevere notizie da nessuna di quelle fonti. Il conteggio totale è la base per le percentuali di colonna; quindi le percentuali di colonna in questo esempio somma a più del 100% visualizzati per la percentuale totale di colonne.

Totali che Do Add Up

Mentre "il conteggio" è tipicamente un termine abbastanza ambiguo, l'esempio che precede dimostra come potrebbe essere confuso nel contesto dei totali per più set di risposta, per i quali *risposte* è spesso la statistica di riepilogo che si desidera davvero.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Fare clic con il tasto destro del mouse su *Origini notizie* nell'anteprima della tabella sul riquadro canvas e selezionare **Statistiche di riepilogo** dal menu a comparsa.
3. Nella finestra di dialogo Statistiche di riepilogo, selezionare **Risposte** nell'elenco Statistiche e clicca sulla freccia per aggiungendola alla lista Visualizza.
4. Selezionare **Risposte Colonna%** nell'elenco Statistiche e fare clic sulla freccia per aggiungendola alla lista Visualizza.
5. Fare clic su **Applica alla selezione**, quindi fare clic su **OK** per creare la tabella.

		Count	Column N %	Responses	Column Responses %
News sources	Get news from internet	867	41.7%	867	24.1%
	Get news from radio	551	26.5%	551	15.3%
	Get news from television	1077	51.8%	1077	30.0%
	Get news from news magazines	294	14.1%	294	8.2%
	Get news from newspapers	805	38.7%	805	22.4%
	Total	2081	100.0%	3594	100.0%

Figura 67. Risposte dicotomie multiple e percentuali di risposta della colonna

Per ogni "categoria" nel set di dicotomia multipla, *Risposte* è identica a *Conteggio*-- e questo sarà sempre il caso per più serie di dicotomia. I totali, però, sono molto diversi. Il numero totale di risposte è pari a 3.594 - oltre il 1.500 in più rispetto al conteggio totale e oltre il 700 in più rispetto al totale dei casi nel file dei dati.

Per le percentuali, i totali per *Colonna N%* e *Risposte Colonne%* sono entrambi al 100% - ma le percentuali per ciascuna categoria nel set di dicotomia multipla sono molto più basse per le percentuali di risposta delle colonne. Questo perché la base percentuale per le percentuali di risposta delle colonne è il numero totale delle risposte, che in questo caso è pari a 3.594, determinando percentuali molto più basse rispetto alla base percentuale di colonna del 2.081.

Percentuale Totali Maggiori Di 100%

Sia le percentuali di colonna che le percentuali di risposta delle colonne restituiscono percentuali totali di 100% anche se, nel nostro esempio, i singoli valori nella colonna *Colonna N%* ammontano chiaramente a un valore superiore al 100%. Quindi, e se volete mostrare percentuali in base al conteggio totale piuttosto che risposte totali ma anche volere la percentuale "totale" per riflettere con precisione la somma delle singole percentuali di categoria?

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Fare clic con il tasto destro del mouse su *Origini notizie* nell'anteprima della tabella sul riquadro canvas e selezionare **Statistiche di riepilogo** dal menu a comparsa.
3. Nella finestra di dialogo Statistiche di riepilogo, selezionare **Risposte Colonna% (Base: Conteggio)** nell'elenco Statistiche e clicca sulla freccia per aggiungendola alla lista Visualizza.
4. Fare clic su **Applica alla selezione**, quindi fare clic su **OK** per creare la tabella.

		Count	Column N %	Responses	Column Responses %	Column Responses % (Base: Count)
News sources	Get news from internet	867	41.7%	867	24.1%	41.7%
	Get news from radio	551	26.5%	551	15.3%	26.5%
	Get news from television	1077	51.8%	1077	30.0%	51.8%
	Get news from news magazines	294	14.1%	294	8.2%	14.1%
	Get news from newspapers	805	38.7%	805	22.4%	38.7%
	Total	2081	100.0%	3594	100.0%	172.7%

Figura 68. Percentuali di risposta della colonna con conteggio come base percentuale

Utilizzo di Più Set di risposta con Altre Variabili

In generale, è possibile utilizzare più serie di risposta esattamente come variabili categoriali. Ad esempio, è possibile crostare una serie di risposta multipla con una variabile categoriale o nidificare una serie di risposta multipla all'interno di una variabile categoriale.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Trascinare e rilasciare *Gender* dall'elenco delle variabili al lato sinistro dell'area Rows sul riquadro in anteprima, nidificando la serie di risposta multipla *News sorgenti* all'interno delle categorie di genere.
3. Fare clic con il tasto destro del mouse su *Gender* nella tabella in anteprima sul riquadro canvas e deselezionare **Mostra Label variabile** sul menu a comparsa.
4. Fare lo stesso per *Fonti News*.

Questo rimuoverà le colonne con le etichette variabili dalla tabella (dato che in questo caso non sono realmente necessarie).

5. Fare clic su **OK** per creare la tabella.

		Count	Column N %	Responses	Column Responses %	Column Responses % (Base: Count)
Male	Get news from internet	359	40.1%	359	23.3%	40.1%
	Get news from radio	233	26.0%	233	15.1%	26.0%
	Get news from television	451	50.3%	451	29.3%	50.3%
	Get news from news magazines	121	13.5%	121	7.9%	13.5%
	Get news from newspapers	375	41.9%	375	24.4%	41.9%
	Total	896	100.0%	1539	100.0%	171.8%
Female	Get news from internet	508	42.9%	508	24.7%	42.9%
	Get news from radio	318	26.8%	318	15.5%	26.8%
	Get news from television	626	52.8%	626	30.5%	52.8%
	Get news from news magazines	173	14.6%	173	8.4%	14.6%
	Get news from newspapers	430	36.3%	430	20.9%	36.3%
	Total	1185	100.0%	2055	100.0%	173.4%

Figura 69. Più serie di risposta nidificate all'interno di una variabile categoriale

Statistiche di riepilogo origine e statistiche di riepilogo disponibili

In assenza di una variabile di scala in una tabella, le variabili categoriali e i set di risposta multipla sono trattati allo stesso modo per quanto riguarda la variabile di origine statistica: La variabile innerpiù nidificata nella dimensione della sorgente di statistiche è la variabile di origine statistica. Dal momento che ci sono alcune statistiche di riepilogo che possono essere assegnate solo a più serie di risposta, ciò significa che il set di risposta multipla deve essere la variabile innerpiù nidificata nella dimensione di origine delle statistiche se si desidera una qualsiasi delle statistiche di riepilogo della risposta multipla.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Nell'anteprima della tabella sul riquadro canvas, trascinare e rilasciare *News sorgenti* a sinistra di *Gender*, modificando l'ordine di nidificazione.

Tutte le statistiche di riepilogo delle risposte multiple speciali -- risposte, percentuali di risposta della colonna -- vengono eliminate dall'anteprima della tabella perché la variabile categoriale *Gender* è ora la variabile innerpiù nidificata e quindi la variabile di origine statistica.

Fortunatamente il costruttore di tabelle "ricorda" queste impostazioni. Se si spostano *News sorgenti* alla sua posizione precedente, nidificate all'interno di *Gender*, tutte le statistiche di riepilogo relative alla risposta vengono ripristinate in anteprima della tabella.

Serie di categorie e risposte duplicate multiple

Più serie di categoria forniscono una funzione non disponibile per i set di dicotomia multipla: la possibilità di contare risposte duplicate. In molti casi, le risposte duplicate in più insiemi di categoria rappresentano probabilmente errori di codifica. Per esempio, per una domanda di indagine come "Quali tre paesi pensi fare le auto migliori?". una risposta di *Svezia, Germania e Svezia* probabilmente non è valida.

In altri casi, tuttavia, le risposte duplicate possono essere perfettamente valide. Ad esempio, se la domanda fosse "Dove sono state le tue ultime tre auto fatte?". una risposta di *Svezia, Germania e Svezia* ha un senso perfetto.

Le tabelle personalizzate forniscono una scelta per le risposte duplicate in più serie di categoria. Per impostazione predefinita, le risposte duplicate non sono conteggiate, ma è possibile richiedere che siano incluse.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Fare clic su **Ripristina** per annullare qualsiasi impostazione precedente.
3. Trascina e rilascia *Car maker, le auto più recenti* dall'elenco delle variabili nell'area Rows del riquadro canvas.

L'icona accanto alla "variabile" nell'elenco delle variabili lo identifica come un insieme di categorie multiple.



Figura 70. Icona serie di più categorie

Per più serie di categorie, le categorie visualizzate rappresentano la serie comune di etichette di valore definite per tutte le variabili nel set (mentre per i set di dicotomia multipla, le "categorie" sono in realtà le etichette variabili per ciascuna variabile nel set).

4. Fare clic con il tasto destro del mouse su *Car maker, la maggior parte delle auto più recenti* nell'anteprima della tabella sul riquadro canvas e selezionare **Categorie e totali** dal menu a comparsa.
5. Selezionare (cliccare) **Totale** nella finestra di dialogo Categorie e Totali, quindi fare clic su **Applica**.
6. Fare clic con il tasto destro del mouse su *Car maker, più recente* e selezionare **Statistiche di riepilogo** dal menu a comparsa.
7. Nella finestra di dialogo Statistiche di riepilogo, selezionare **Risposte** nell'elenco Statistiche e clicca sulla freccia per aggiungendola alla lista Visualizza.
8. Fare clic su **Applica alla selezione**, quindi fare clic su **OK** per creare la tabella.

		Count	Responses
Car maker, most recent cars	American	1938	1938
	Japanese	1327	1327
	Korean	695	695
	German	693	693
	Swedish	360	360
	Other	343	343
Total		2832	5356

Figura 71. Serie multiple di categoria: Conti e risposte senza duplicati

Per impostazione predefinita, le risposte duplicate non sono conteggiate; quindi in questa tabella i valori per ciascuna categoria nelle colonne *Conteggio* e *Risposte* sono identici. Solo i totali differiscono.

9. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).

10. Fare clic sulla scheda **Opzioni**.
11. Clicca (spunta) **Conteggio risposte duplicate per più serie di categorie**.
12. Fare clic su **OK** per creare la tabella.

		Count	Responses
Car maker, most recent cars	American	1938	2797
	Japanese	1327	1717
	Korean	695	760
	German	693	754
	Swedish	360	383
	Other	343	359
Total		2832	6770

Figura 72. Serie multipla con risposte duplicate incluse

In questa tabella c'è una notevole differenza tra i valori nelle colonne *Conteggio* e *Risposte*, in particolare per le auto americane, indicando che molti intervistati hanno posseduto più auto americane.

Significato Testing con Più Set di risposta

È possibile utilizzare più serie di risposta in test di significatività in essenzialmente allo stesso modo in cui si utilizzerebbero variabili categoriali.

- Per i test di indipendenza (chi-square) o confrontare proporzioni di colonna (z - test), i test vengono eseguiti su conteggi e il conteggio deve essere una delle statistiche di riepilogo visualizzate nella tabella.
- Per più serie di categorie, i test che mettono a confronto proporzioni di colonne o mezzi di colonna (t - test) non vengono eseguiti se **Conteggio risposte duplicate per più serie di categorie** viene selezionata nella scheda Opzioni. Per ulteriori informazioni, consultare l'argomento ["Tabelle personalizzate: Opzioni Tab"](#) a pagina 15.

Test di Indipendenza con Più Set di risposta

Questo esempio crea una crosstabulazione di una variabile categoriale e una serie di risposta multipla ed esegue un test di chi - quadrato di indipendenza sulla crosstabulazione.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Fare clic su **Ripristina** per annullare qualsiasi impostazione precedente.
3. Trascinare e rilasciare le *Fonti News* (questa è l'etichetta descrittiva per la serie di dicotomia multipla \$mltnews) dall'elenco delle variabili nell'area Colonne del riquadro canvas.
4. Trascinare e rilasciare *Gender* dall'elenco delle variabili nell'area Rows del riquadro canvas.
5. Fare clic sulla scheda **Statistiche di test**.
6. Selezionare (controllare) **Test di indipendenza (chi-square)**.
7. Se non è già selezionato, selezionare **Includi più variabili di risposta in test**.
8. Fare clic su **OK** per eseguire la procedura.

		News sources
Gender	Chi-square	10.266
	df	5
	Sig.	.068

Figura 73. Risultati di chi - quadrato

Il livello di significatività di 0.068 per il test del chi - quadrato indica che probabilmente maschi e femmine non differiscono in modo significativo nelle scelte delle fonti di notizie (supponendo di utilizzare un valore di significatività di 0.05 o inferiore come criterio per determinare la significatività statistica).

Confronto Colonna Significa con Più Set di risposta

Questo esempio calcola mezzi di una variabile di scala all'interno delle categorie definite da una serie di risposta multipla e confronta ogni categoria media ad ogni altra categoria media per differenze significative.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Fare clic su **Ripristina** per annullare qualsiasi impostazione precedente.
3. Trascinare e rilasciare le *Fonti News* (questa è l'etichetta descrittiva per la serie di dicotomia multipla \$mltnews) dall'elenco delle variabili nell'area Colonne del riquadro canvas.
4. Trascinare e rilasciare *Age of respondent* nell'area Rows del riquadro canvas.
5. Fare clic sulla scheda **Statistiche di test**.
6. Selezionare (controllare) **Confronta Colonna Mezzi (t - test)**.
7. Se non è già selezionato, selezionare **Includi più variabili di risposta in test**.
8. Fare clic su **OK** per eseguire la procedura.

	News sources				
	Get news from newspapers	Get news from news magazines	Get news from television	Get news from radio	Get news from internet
	Mean	Mean	Mean	Mean	Mean
Age of respondent	52	40	48	40	40

Comparisons of Column Means

	News sources				
	Get news from newspapers	Get news from news magazines	Get news from television	Get news from radio	Get news from internet
	(A)	(B)	(C)	(D)	(E)
Age of respondent	B C D E		B D E		

Results are based on two-sided tests assuming equal variances with significance level 0.05. For each significant pair, the key of the smaller category appears under the category with larger mean.

Figura 74. Risultati del test di significato

- Ogni categoria del set di risposta multipla è identificata da una lettera (A, B, C, D, E) e per ciascuna categoria per cui la media di un'altra categoria è inferiore e differisce sensibilmente dalla media di quella categoria, viene visualizzata la lettera che rappresenta la categoria con la media inferiore.
- *Ottieni notizie dai giornali* (A) ha l'età media più alta, e tutte le altre categorie significa differire in modo significativo da esso.
- *Ottieni notizie dalla televisione* (C) ha la prossima età media più alta, e tutti i restanti mezzi di categoria (B, D ed E) differiscono in modo significativo da esso. (C anche differisce sensibilmente da A, come precedentemente indicato.)
- Le età medie per *Ottieni notizie dalla rivista* (B), *Ottieni notizie dalla radio* (D) e *Ottieni notizie da internet* (E) non differiscono significativamente l'una dall'altra.

Valori mancanti

Molti file di dati contengono una certa quantità di dati mancanti. Un'ampia varietà di fattori può causare i dati mancanti. Ad esempio, gli intervistati intervistati possono non rispondere ad ogni domanda, alcune variabili potrebbero non essere applicabili ad alcuni casi, e gli errori di codifica possono determinare alcuni valori buttati fuori.

Ci sono due tipi di valori mancanti in IBM SPSS Statistica :

- **Utente - mancante.** Valori definiti come contenenti dati mancanti. Le etichette di valore possono essere assegnate a questi valori per identificare il motivo per cui mancano i dati (come ad esempio un codice di 99 e un'etichetta di valore di *Non Applicabile* per la gravidanza nei maschi).
- **Sistema - mancante.** Se nessun valore è presente per una variabile numerica, viene assegnato il valore mancante di sistema. Questo è indicato da un periodo nella Data View dell'Editor dei dati.

Ci sono una serie di strutture che possono aiutare a compensare gli effetti dei dati mancanti e anche analizzare i modelli nei dati mancanti. Questo capitolo, tuttavia, ha un obiettivo molto più semplice: descrivere come le tabelle personalizzate gestiscono i dati mancanti e come i dati mancanti influenzano il calcolo delle statistiche di riepilogo.

File di dati di esempio

Gli esempi in questo capitolo utilizzano il file di dati *missing_values.sav*. Per ulteriori informazioni, consultare l'argomento ["File di esempio"](#) a pagina 76. Si tratta di un file di dati molto semplice, completamente artificiale, con una sola variabile e dieci casi, studiati per illustrare i concetti di base sui valori mancanti.

Tabelle senza valori mancanti

Per impostazione predefinita, le categorie mancanti dell'utente non vengono visualizzate nelle tabelle personalizzate (e i valori mancanti di sistema non vengono mai visualizzati).

1. Dai menu, scegliere:

Analisi > Tavole > Tabelle personalizzate ...

2. Nel builder table, trascinare e rilasciare *Variabile con valori mancanti* (l'unica variabile nel file) dall'elenco variabile nell'area Rows del riquadro canvas.

3. Fare clic con il tasto destro del mouse sulla variabile sul riquadro canvas e selezionare **Categorie e Totali** dal menu a comparsa.

4. Fare clic (controllare) **Totale** nella finestra di dialogo Categorie e Totali, quindi fare clic su **Applica**.

5. Fare clic con il tasto destro del mouse su *Variabile con i valori mancanti* nella tabella in anteprima sul riquadro canvas e selezionare **Statistiche riepilogo** dal menu a comparsa.

6. Nella finestra di dialogo Statistiche di riepilogo, selezionare **Colonna N%** nell'elenco Statistiche e fare clic sulla freccia per aggiungendola alla lista Visualizza.

7. Fare clic su **Applica alla selezione**.

Si può notare una leggera discrepanza tra le categorie visualizzate nell'anteprima della tabella sul riquadro canvas e le categorie visualizzate nell'elenco Categorie (sotto l'elenco delle variabili sul lato sinistro del costruttore di tabelle). L'Elenco Categorie contiene una categoria etichettata *Valori mancanti* che non è inclusa nell'anteprima della tabella perché le categorie di valore mancanti sono escluse per impostazione predefinita. Dal momento che i "valori" sono plurali in etichetta, questo indica che la variabile ha due o più categorie mancanti di utenti.

8. Fare clic su **OK** per creare la tabella.

		Count	Column N %
Variable with missing values	Low	2	28.6%
	Medium	3	42.9%
	High	2	28.6%
	Total	7	100.0%

Figura 75. Tabella senza valori mancanti

Tutto in questo tavolo va benissimo. I valori di categoria vengono sommati ai totali e le percentuali riflettono accuratamente i valori che si otterrebbero utilizzando il conteggio totale come base percentuale (ad esempio, $3/7 = 0.4290$ 42.9%). Il conteggio totale, tuttavia, non è il numero totale di casi nel file dei dati; è il numero totale di casi con valori **non mancanti** o i casi che non hanno valori mancanti di utente o di sistema per quella variabile.

Inclusi i valori mancanti nelle tabelle

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Fare clic con il tasto destro del mouse su *Variabile con valori mancanti* nell'anteprima della tabella sul riquadro canvas e selezionare **Categorie e Totali** dal menu a comparsa.
3. Clicca (verificare) **Valori mancanti** nella finestra di dialogo Categorie e Totali, quindi fare clic su **Applica**.

Ora l'anteprima della tabella include una categoria *Valori mancanti*. Anche se l'anteprima della tabella visualizza solo una categoria per i valori mancanti, tutte le categorie mancanti verranno visualizzate nella tabella.

4. Fare clic con il tasto destro del mouse su *Variabile con i valori mancanti* nella tabella in anteprima sul riquadro canvas e selezionare **Statistiche riepilogo** dal menu a comparsa.
5. Nella finestra di dialogo Statistiche di riepilogo, clicca (spunta) **Statistiche riepilogo personalizzate per Totali e Subtotali**.
6. Selezionare **N valido** nell'elenco Statistiche di riepilogo personalizzate e fare clic sulla freccia per aggiungendola alla lista Visualizza.
7. Fare lo stesso per **Total N**.
8. Fare clic su **Applica alla selezione**, quindi fare clic su **OK** nel builder da tavolo per creare la tabella.

		Count	Column N %	Valid N	Total N
Variable with missing values	Low	2	22.2%		
	Medium	3	33.3%		
	High	2	22.2%		
	Don't know	1	11.1%		
	Not applicable	1	11.1%		
	Total	9	100.0%		
				7	10

Figura 76. Tabella con valori mancanti

Le due categorie definite dall'utente definite *-Non conoscere e Non applicabile-* sono ora visualizzate nella tabella e il conteggio totale è ora di 9 invece di 7, riflettendo l'aggiunta dei due casi con valori mancanti dell'utente (uno in ogni categoria di utenti - mancante). Anche le percentuali di colonna sono diverse ora, perché si basano sul numero di valori non mancanti e mancanti degli utenti. Solo i valori mancanti di sistema non sono inclusi nel calcolo percentuale.

Valido N mostra il numero totale di casi non mancanti (7) e *Total N* mostra il numero totale di casi, inclusi entrambi gli utenti mancanti e quelli di sistema. Il numero totale dei casi è di 10, uno in più rispetto al conteggio dei valori non mancanti e quelli mancanti degli utenti visualizzati come il totale nella colonna *Conteggio*. Questo perché c'è un caso con un valore mancante di sistema.

9. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
10. Fare clic con il tasto destro del mouse su *Variabile con valori mancanti* nell'anteprima della tabella sul riquadro canvas e selezionare **Statistiche di riepilogo** dal menu a comparsa.
11. Selezionare **Colonna valida N%** nell'elenco delle principali statistiche (non i riepiloghi personalizzati per i totali e i subtotali) e clicca sulla freccia per aggiungerlo all'elenco Visualizza.
12. Fare lo stesso per **Colonna Total N%**.
13. È anche possibile aggiungerli sia all'elenco delle statistiche di riepilogo personalizzate per i totali e i subtotali.
14. Fare clic su **Applica alla selezione**, quindi fare clic su **OK** per creare la tabella.

		Count	Column N %	Column Valid N %	Column Total N %	Valid N	Total N
Variable with missing values	Low	2	22.2%	28.6%	20.0%		
	Medium	3	33.3%	42.9%	30.0%		
	High	2	22.2%	28.6%	20.0%		
	Don't know	1	11.1%	.0%	10.0%		
	Not applicable	1	11.1%	.0%	10.0%		
	Total	9	100.0%	100.0%	100.0%	7	10

Figura 77. Tabella con valori mancanti e percentuali valide e totali

- **Colonna N%** è la percentuale in ogni categoria in base al numero di valori mancanti e mancanti di utenti (dato che i valori mancanti degli utenti sono stati esplicitamente inclusi nella tabella).
- **Colonna Valido N%** è la percentuale in ogni categoria in base ai soli casi validi, non mancanti. Questi valori sono gli stessi delle percentuali di colonna che si trovavano nella tabella originale che non include i valori mancanti dell'utente.
- **Colonna totale N%** è la percentuale in ogni categoria basata su tutti i casi, inclusi sia l'utente - mancante che il sistema - mancante. Se si sommano le percentuali di singole categorie in questa categoria, si vedrà che si aggiungono fino a soli 90%, perché un caso su totale di 10 casi (10%) ha il valore mancante di sistema. Sebbene questo caso sia incluso nella base per i calcoli percentuali, nella tabella non viene fornita alcuna categoria per i casi con valori mancanti di sistema.

Formattazione e personalizzazione delle tabelle

Formattazione e personalizzazione delle tabelle

Le Tabelle personalizzate forniscono la possibilità di controllare una serie di proprietà di formattazione della tabella come parte del processo di costruzione della tabella, tra cui:

- Formato di visualizzazione e etichette per le statistiche di riepilogo
- Larghezza minima e massima della colonna dei dati
- Testo o valore visualizzato nelle celle vuote

Queste impostazioni persistono all'interno dell'interfaccia del builder table (fino a modificarle, reimpostare le impostazioni del builder della tabella o aprire un diverso file di dati), consentendo di creare più tabelle con le stesse proprietà di formattazione senza modificare manualmente le tabelle dopo la loro creazione. È anche possibile salvare queste impostazioni di formattazione, insieme a tutti gli altri parametri di tabella, utilizzando il pulsante Paste nell'interfaccia del builder di tabella per incollare la sintassi dei comandi in una finestra di sintassi, che è possibile quindi salvare come file.

È anche possibile modificare molte proprietà di formattazione delle tabelle dopo che sono state create, utilizzando tutte le funzionalità di formattazione disponibili nel Viewer per le tabelle pivot. Questo capitolo, tuttavia, si concentra sul controllo delle proprietà di formattazione delle tabelle prima della creazione della tabella. Per ulteriori informazioni sulle tabelle pivot, utilizzare la scheda Indice nel sistema della Guida e immettere `pivot tables` come parola chiave.

File di dati di esempio

Gli esempi in questo capitolo utilizzano il file di dati *survey_sample.sav*. Per ulteriori informazioni, consultare l'argomento [“File di esempio”](#) a pagina 76.

Tutti gli esempi forniti qui visualizzano etichette variabili in finestre di dialogo, ordinate in ordine alfabetico. Le proprietà di visualizzazione della lista variabili sono impostate sulla scheda Generale nella finestra di dialogo Opzioni (Modifica menu, Opzioni).

Formato visualizzazione statistiche di riepilogo

Le tabelle personalizzate tentano di applicare formati predefiniti relativamente intelligenti alle statistiche di riepilogo, ma probabilmente ci saranno momenti in cui si desidera sovrascrivere questi default.

1. Dai menu, scegliere:
Analisi > Tavole > Tabelle personalizzate ...
2. Nel builder table, trascinare e rilasciare la *Categoria Età* dall'elenco delle variabili nell'area Rows sul riquadro canvas.
3. Trascinare e rilasciare *Confidence in televisione* sotto *Categoria Età* nell'area Rows, impilando le due variabili nella dimensione della riga.
4. Fare clic con il tasto destro del mouse su *Categoria di età* nella tabella in anteprima sul riquadro canvas e selezionare **Seleziona tutte le variabili di riga** dal menu a comparsa.
5. Fare clic con il tasto destro del mouse su *Categoria di età* e selezionare **Categorie e totali** dal menu a comparsa.
6. Nella finestra di dialogo Categorie e Totali, selezionare (controllare) **Totale** e quindi fare clic su **Applica**.
7. Fare clic con il tasto destro del mouse sulla variabile in anteprima sul riquadro canvas e selezionare **Statistiche riepilogo** dal menu a comparsa.
8. Selezionare **Colonna N%** nell'elenco Statistiche e fare clic sul tasto freccia per aggiungerlo all'elenco Visualizza.
9. Selezionare (controllare) **Statistiche riepilogo personalizzate per Totali e Subtotali**.
10. Nell'elenco Statistiche per le statistiche di riepilogo personalizzate, selezionare **Colonna N%** e clicca sulla freccia per aggiungendola alla lista Visualizza.
11. Fare lo stesso per **Mean**.
12. Quindi fare clic su **Applica a tutti**.

I valori segnaposto nell'anteprima della tabella riflettono il formato predefinito per ogni statistica di riepilogo.

- Per i conteggi, il formato di visualizzazione predefinito è **nnnn--** valori interi senza decimali.
- Per le percentuali, il formato di visualizzazione predefinito è **nnnn.n%** numeri con una singola posizione decimale e un segno percentuale dopo il valore.
- Per la media, il formato di visualizzazione predefinito è *diverso* per le due variabili.

Per le statistiche di riepilogo che non sono una qualche forma di conteggio (inclusi N e Total N) o percentuale, il formato di visualizzazione predefinito è il formato di visualizzazione definito per la variabile nell'Editor dei dati. Se si osservano le variabili in Variabile Vista nell'Editor dei dati, si vedrà che *Categoria di età* (variabile *agecat*) è definita come avere due posizioni decimali, mentre *Confidence in televisione* (variabile *contv*) è definita come avere zero posizioni decimali.

Questo è uno di quei casi in cui il formato predefinito probabilmente non è il formato desiderato, dato che probabilmente sarebbe meglio se entrambi i valori medi visualizzavano lo stesso numero di decimali.

13. Fare clic con il tasto destro del mouse sulla variabile in anteprima sul riquadro canvas e selezionare **Statistiche riepilogo** dal menu a comparsa.
 Per la media, la cella Formato nell'elenco Visualizza indica che il formato è *Auto*, il che significa che verrà utilizzato il formato di visualizzazione definito per la variabile e la cella Decimali è disabilitata. Per specificare il numero di decimali, è necessario innanzitutto selezionare un formato diverso.
14. Nell'elenco delle statistiche di riepilogo personalizzate, fare clic sulla cella Formato per la media e selezionare **nnnn** dall'elenco a discesa dei formati.
15. Nella cella Decimali, immettere il valore 1.
16. Quindi fare clic su **Applica a tutti** per applicare questa impostazione a entrambe le variabili.

Ora l'anteprima della tabella indica che entrambi i valori medi verranno visualizzati con una posizione decimale. (Si potrebbe andare avanti e creare questo tavolo ora - ma potreste trovare il valore "medio" per *Categoria Età* un po' difficile da interpretare, dato che i codici numerici effettivi per questa variabile variano solo da 1 a 6.)

Visualizza Etichette per statistiche di riepilogo

Oltre ai formati di visualizzazione per le statistiche di riepilogo, è possibile anche controllare le etichette descrittive per ogni statistica di riepilogo.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
 2. Clicca su **Reset** per cancellare le eventuali impostazioni precedenti nel builder table.
 3. Nel builder table, trascinare e rilasciare la *Categoria Età* dall'elenco delle variabili nell'area Rows sul riquadro canvas.
 4. Trascina e rilascia *Come viene pagato la scorsa settimana* dall'elenco delle variabili nell'area Colonne sul riquadro canvas.
 5. Fare clic con il tasto destro del mouse su *Categoria di età* nella tabella in anteprima sul riquadro canvas e selezionare **Statistiche di riepilogo** dal menu del contesto pop - up.
 6. Selezionare **Colonna N%** nell'elenco Statistiche e fare clic sul tasto freccia per aggiungerlo all'elenco Visualizza.
 7. Fare doppio clic ovunque nella parola *Colonna* nella cella Label nell'elenco Visualizza per modificare il contenuto della cella. Eliminare la parola *Colonna* dall'etichetta, modificando l'etichetta in modo semplice %.
 8. Modificare la cella Label per *Conteggio* nello stesso modo, modificando l'etichetta in modo semplice N.
- Mentre siamo qui, modifichiamo il formato della statistica N% per rimuovere l'inutile segno percentuale (dato che l'etichetta della colonna indica che la colonna contiene percentuali).
9. Fare clic sulla cella Formato per % N colonna e selezionare **nnnn.n** dall'elenco a discesa dei formati.
 10. Quindi fare clic su **Applica alla selezione**.

L'anteprima della tabella visualizza il formato di visualizzazione modificato e le etichette modificate.

11. Fare clic su **OK** per creare la tabella.

		How get paid last week											
		Hourly wage		Daily wage		Weekly wage		Monthly salary		Annual salary		Other pay rate	
		N	%	N	%	N	%	N	%	N	%	N	%
Age category	Less than 25	91	14.0	0	.0	12	9.7	3	2.0	7	3.1	14	7.7
	25 to 34	175	26.9	5	29.4	33	26.6	37	24.8	63	28.0	31	17.1
	35 to 44	185	28.5	5	29.4	42	33.9	45	30.2	66	29.3	61	33.7
	45 to 54	124	19.1	5	29.4	25	20.2	38	25.5	58	25.8	41	22.7
	55 to 64	52	8.0	0	.0	10	8.1	23	15.4	29	12.9	19	10.5
	65 or older	23	3.5	2	11.8	2	1.6	3	2.0	2	.9	15	8.3

Figura 78. Tabella con etichette statistiche di riepilogo modificate

Larghezza colonna

Avrete forse notato che la tabella nell'esempio precedente è piuttosto ampia. Una soluzione a questo problema sarebbe quella di scambiare semplicemente le variabili di riga e colonna. Un'altra soluzione è quella di rendere le colonne più strette, dato che sembrano essere molto più ampie del necessario. (In effetti, il motivo per cui abbiamo accorciato le etichette statistiche di riepilogo è stato così da rendere le colonne più strette).

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Fare clic sulla scheda **Opzioni**.
3. Nel gruppo Width for Data Columns, selezionare **Custom**.
4. Per il valore massimo, immettere 36. (Assicurarsi che l'impostazione Unità sia **Punti**.)
5. Fare clic su **OK** per creare la tabella.

		How get paid last week											
		Hourly wage		Daily wage		Weekly wage		Monthly salary		Annual salary		Other pay rate	
		N	%	N	%	N	%	N	%	N	%	N	%
Age category	Less than 25	91	14.0	0	.0	12	9.7	3	2.0	7	3.1	14	7.7
	25 to 34	175	26.9	5	29.4	33	26.6	37	24.8	63	28.0	31	17.1
	35 to 44	185	28.5	5	29.4	42	33.9	45	30.2	66	29.3	61	33.7
	45 to 54	124	19.1	5	29.4	25	20.2	38	25.5	58	25.8	41	22.7
	55 to 64	52	8.0	0	.0	10	8.1	23	15.4	29	12.9	19	10.5
	65 or older	23	3.5	2	11.8	2	1.6	3	2.0	2	.9	15	8.3

Figura 79. Tabella con larghezze di colonna ridotte

Ora il tavolo è molto più compatto.

Valore di visualizzazione per celle vuote

Per impostazione predefinita, viene visualizzato un 0 in celle vuote (celle che non contengono casi). È possibile invece visualizzare nulla in queste celle (lasciarle in bianco) oppure specificare una stringa di testo da visualizzare in celle vuote.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Fare clic sulla scheda **Opzioni**.
3. Nel gruppo Aspetto cella dati, per Celle vuote selezionare **Testo** e immettere None.
4. Fare clic su **OK** per creare la tabella.

		How get paid last week											
		Hourly wage		Daily wage		Weekly wage		Monthly salary		Annual salary		Other pay rate	
		N	%	N	%	N	%	N	%	N	%	N	%
Age category	Less than 25	91	14.0	None	None	12	9.7	3	2.0	7	3.1	14	7.7
	25 to 34	175	26.9	5	29.4	33	26.6	37	24.8	63	28.0	31	17.1
	35 to 44	185	28.5	5	29.4	42	33.9	45	30.2	66	29.3	61	33.7
	45 to 54	124	19.1	5	29.4	25	20.2	38	25.5	58	25.8	41	22.7
	55 to 64	52	8.0	None	None	10	8.1	23	15.4	29	12.9	19	10.5
	65 or older	23	3.5	2	11.8	2	1.6	3	2.0	2	.9	15	8.3

Figura 80. Tabella con "Nessuno" visualizzato nelle celle vuote

Ora le quattro celle vuote nella tabella visualizzano il testo *Nessuno* invece di un valore di 0.

Valore di visualizzazione per statistiche mancanti

Se una statistica non può essere calcolata, il valore di visualizzazione predefinito è un periodo (.), ovvero il simbolo utilizzato per indicare il valore mancante di sistema. Questo è diverso da una cella "vuota" e quindi il valore di visualizzazione per le statistiche mancanti è controllato separatamente dal valore di visualizzazione per le celle che non contengono casi.

1. Aprire il builder table (Analizza menu, tabelle, Tabelle personalizzate).
2. Trascinare e rilasciare *Ore al giorno guardando la TV* dall'elenco delle variabili alla parte superiore dell'area delle Colonne sui canvas, sopra *Come si pagano la scorsa settimana*.

Dal momento che *Ore al giorno guardando la TV* è una variabile di scala, diventa automaticamente la variabile di origine statistica e la statistica di riepilogo cambia alla media.

3. Fare clic con il tasto destro del mouse su *Ore al giorno guardando la TV* nell'anteprima della tabella nel riquadro canvas e selezionare **Statistiche di riepilogo** dal menu del contesto pop-up.
4. Selezionare **N valido** nell'elenco Statistiche e fare clic sul tasto freccia per aggiungerlo all'elenco Visualizza.
5. Fare clic su **Applica alla selezione**.
6. Fare clic sulla scheda **Opzioni**.
7. Nel campo di testo per Statistiche che non possono essere calcolate, immettere NA.
8. Fare clic su **OK** per creare la tabella.

		Hours per day watching TV											
		How get paid last week											
		Hourly wage		Daily wage		Weekly wage		Monthly salary		Annual salary		Other pay rate	
		Mean	Valid N	Mean	Valid N	Mean	Valid N	Mean	Valid N	Mean	Valid N	Mean	Valid N
Age category	Less than 25	3	71	NA	None	3	10	2	3	2	6	2	8
	25 to 34	3	134	5	2	2	30	2	29	2	52	2	22
	35 to 44	3	136	2	5	3	30	2	34	2	47	3	46
	45 to 54	2	90	2	4	2	22	2	36	2	45	2	34
	55 to 64	3	40	NA	None	3	7	2	15	2	23	3	15
	65 or older	3	18	2	2	1	1	NA	0	1	2	3	11

Figura 81. Tabella con "NA" visualizzato per le statistiche mancanti

Il testo NA viene visualizzato per la media in tre celle della tabella. In ogni caso, il valore corrispondente Valido N spiega perché: Non ci sono casi con cui calcolare la media.

Si può, tuttavia, notare quello che sembra essere una leggera discrepanza - uno di quei tre valori N validi viene visualizzato come un 0, piuttosto che l'etichetta *Nessuno* che dovrebbe essere visualizzata in celle senza casi. Questo perché sebbene non ci siano casi validi da usare per calcolare la media, la categoria non è davvero vuota. Se si torna al tavolo originale con solo le due variabili categoriali, si vedrà che ci sono, infatti, tre casi in questa categoria crosstabulata. Non ci sono casi validi, tuttavia, perché tutti e tre hanno valori mancanti per la variabile di scala *Ore al giorno guardando la TV*.

File di esempio

Il file di esempio installato con il prodotto si trova nella sottocartella *Samples* della cartella di installazione. All'interno della sottodirectory *Samples* è presente una cartella separata per ciascuna delle seguenti lingue: Inglese, francese, tedesco, italiano, giapponese, coreano, polacco, , cinese semplificato, spagnolo e cinese tradizionale.

Non tutti i file di esempio sono disponibili in tutte le lingue. Se un file di esempio non è disponibile in una lingua, la cartella di tale lingua contiene una versione inglese del file.

Descrizioni

Questa sezione contiene brevi descrizioni dei file di esempio utilizzati negli esempi riportati in tutta la documentazione.

- **accidents.sav.** File di dati ipotetici che prende in esame una compagnia di assicurazioni impegnata nello studio dei fattori di rischio correlati all'età e al sesso per gli incidenti automobilistici che si verificano in una determinata regione. Ciascun caso corrisponde a una classificazione incrociata della categoria relativa età e del sesso.
- **adl.sav.** File di dati ipotetici che prende in esame l'impegno richiesto per determinare i vantaggi di un tipo di terapia proposto per i pazienti con problemi di cuore. I medici hanno assegnato in modo casuale i pazienti con problemi di cuore di sesso femminile a uno di due gruppi. Al primo gruppo è stata assegnata la terapia fisica standard; al secondo gruppo, un'ulteriore terapia di supporto psicologico. Dopo tre mesi di trattamenti, a ciascuna capacità dei pazienti che consente di riprendere le normali attività giornaliere è stato assegnato un punteggio come variabile ordinale.
- **advert.sav.** File di dati ipotetici che prende in esame l'impegno di un rivenditore al dettaglio che desidera esaminare la relazione tra il denaro speso per la pubblicità e le vendite risultanti. Finora sono stati raccolti i dati delle vendite precedenti e i relativi costi pubblicitari.
- **aflatoxin.sav.** File di dati ipotetici che prende in esame il test di raccolti di mais con presenza di Aflatossina, un veleno la cui concentrazione varia notevolmente nei raccolti. Una macchina per la lavorazione dei cereali ha ricevuto 16 campioni da ciascuno degli otto raccolti di mais e ha misurato i livelli di Aflatossina in parti per miliardo (PPB).
- **anorectic.sav.** Mentre si lavora verso una sintomatologia standardizzata del comportamento anorectico / bulimico, i ricercatori¹ fatta uno studio su 55 adolescenti con disturbi alimentari noti. Ogni paziente è stato visitato quattro volte in quattro anni, per un totale di 220 visite. Durante ogni visita, ai pazienti sono stati assegnati punteggi per ciascuno dei 16 sintomi. I punteggi relativi ai sintomi sono

¹ Van der Ham, T., J. J. Meulman, D. C. Van Strien e H. Van Engeland. 1997. Empirically based subgrouping of eating disorders in adolescents: A longitudinal perspective. *British Journal of Psychiatry*, 170, 363-368.

assenti per il paziente 71 alla visita 2, il paziente 76 alla visita 2 e il paziente 47 alla visita 3, con 217 osservazioni valide.

- **bankloan.sav.** File di dati ipotetici che prende in esame l'impegno di una banca nel tentativo di ridurre il tasso di inadempienza nel rimborso di un prestito. Il file contiene informazioni finanziarie e demografiche su 850 vecchi e potenziali clienti. I primi 700 casi riguardano i clienti a cui sono stati concessi dei prestiti precedentemente. Gli ultimi 150 casi riguardano i potenziali clienti che la banca deve classificare come rischi di credito positivi o negativi.
- **bankloan_binning.sav.** File di dati ipotetici che contiene informazioni finanziarie e demografiche su 5000 vecchi clienti.
- **behavior.sav.** In un esempio classico² 52 studenti è stato chiesto di valutare le combinazioni di 15 situazioni e 15 comportamenti su una scala di 10 punti che varia da 0 = "estremamente appropriata" a 9 = "estremamente inappropriato". I valori medi riferiti ai partecipanti sono stati considerati dissimilarità.
- **behavior_ini.sav.** Questo file di dati contiene la configurazione iniziale di una soluzione a due dimensioni per *behavior.sav*.
- **brakes.sav.** File di dati ipotetici che prende in esame il controllo di qualità di un'industria che produce freni a disco per automobili con elevata prestazione. Il file di dati contiene le misurazioni del diametro di 16 dischi da ciascuna delle otto macchine di produzione. L'obiettivo finale è ottenere un diametro dei dischi pari a 322 millimetri.
- **breakfast.sav.** In uno studio classico³, 21 studenti di Wharton School MBA e i loro coniugi sono stati chiamati a classificare 15 articoli per la colazione in ordine di preferenza con 1 = "più preferite" a 15 = "meno preferite". Le loro preferenze sono state registrate per sei diversi scenari, che comprendevano tutti gli scenari compresi tra "Preferenza generale" e "Solo snack con bibita".
- **breakfast-overall.sav.** Questo file contiene le preferenze degli alimenti della colazione solo per il primo scenario, "Preferenza generale".
- **broadband_1.sav.** File di dati ipotetici che contiene il numero di sottoscrittori, per area, di un provider di servizi a banda larga nazionale. Il file di dati contiene il numero dei sottoscrittori mensili di 85 aree in un periodo di quattro anni.
- **broadband_2.sav.** Questo file è identico al file *broadband_1.sav*, ma contiene i dati per ulteriori tre mesi.
- **car_insurance_claims.sav.** Un dataset presentato e analizzato altrove⁴ riguarda i danni per le auto. La quantità media di richieste di risarcimento può essere adattata come avente una distribuzione gamma, utilizzando una funzione di collegamento inverso per correlare la media della variabile dipendente a una combinazione lineare di età del contraente della polizza e tipo e anni del veicolo. Il numero delle richieste di risarcimento specificato può essere utilizzato come peso di scaling.
- **car_sales.sav.** Questo file di dati ipotetici contiene le stime sulle vendite, i prezzi di listino e le specifiche fisiche di numerose marche e modelli di veicoli. I prezzi di listino e le specifiche fisiche sono state ottenute dal sito *edmunds.com* e dai siti dei produttori.
- **car_sales_uprepared.sav.** Questa è una versione modificata di *car_sales.sav* che non include le versioni trasformate dei campi.
- **carpet.sav.** In un esempio popolare⁵, una società interessata a commercializzare un nuovo detergente per i tappeti vuole esaminare l'influenza di cinque fattori sulla preferenza del consumatore - pacchetto design, marchio, prezzo, un sigillo *Good Housekeeping* e una garanzia money-back. Esistono tre livelli di fattori per il disegno del package, che differiscono per l'ubicazione della spazzola dell'applicatore; tre marchi (*K2R*, *Glory* e *Bissell*); tre livelli di prezzo e due livelli (no o sì) per ciascuno degli ultimi due fattori. Dieci consumatori sono classificati in 22 profili definiti da questi fattori. La variabile *Preferenza*

² Prezzo, R. H., e D. L. Bouffard. 1974. Behavioral appropriateness and situational constraints as dimensions of social behavior. *Journal of Personality and Social Psychology*, 30, 579-586.

³ Verde, P. E., e V. Rao. 1972. *Applied multidimensional scaling*. Hinsdale, Ill.: Dryden Press.

⁴ McCullagh, P., e J. A. Nelder. 1989. *Modelli lineari generalizzati*, 2a ed. Londra: Chapman & Hall.

⁵ Verde, P. E., e Y. Vento. 1973. *Multiattribute decisions in marketing: A measurement approach*. Hinsdale, Ill.: Dryden Press.

contiene come classificazione il rango medio per ogni profilo. Classificazioni basse corrispondono a una preferenza elevata. La variabile riflette una misura globale della preferenza per ogni profilo.

- **carpet_prefs.sav.** Questo file di dati si basa sullo stesso esempio del file *carpet.sav*, ma contiene le classificazioni effettive raccolte da ciascuno dei 10 clienti. Ai clienti è stato chiesto di classificare 22 profili di prodotti in ordine di preferenza. Le variabili da *PREF1* a *PREF22* contengono gli ID dei profili associati, come definito nel file *carpet_plan.sav*.
- **catalog.sav.** File di dati ipotetico che contiene le cifre sulle vendite mensili di tre prodotti venduti da una società di vendita per corrispondenza. Il file include anche i dati di cinque possibili variabili predittore.
- **catalog_seasfac.sav.** Questo file di dati è uguale al file *catalog.sav* con l'eccezione che contiene un insieme di fattori stagionali calcolati dalla procedura Decomposizionale stagionale insieme a variabili di dati.
- **cellular.sav.** File di dati ipotetici che prende in esame l'impegno di un'azienda di telefonia cellulare nel tentativo di ridurre il tasso di abbandono, ovvero l'abbandono dei clienti. Vengono applicati agli account dei punteggi di propensione all'abbandono, che vanno da 0 a 100. Gli account che totalizzano 50 o più potrebbero essere propensi a cambiare fornitore.
- **ceramics.sav.** File di dati ipotetici che prende in esame l'impegno di un produttore che desidera stabilire se una nuova lega premium ha una maggiore resistenza al calore rispetto alla lega standard. Ciascun caso rappresenta il test separato di una delle leghe. È indicata la temperatura massima alla quale può essere sottoposto il cuscinetto.
- **cereal.sav.** File di dati ipotetici che prende in esame le preferenze relative agli alimenti della colazione di un campione di 880 persone. Il file riporta anche l'età, il sesso e lo stato civile del campione e se le persone conducono uno stile di vita attivo (in base a un'attività sportiva con frequenza di due volte alla settimana). Ogni caso rappresenta un rispondente separato.
- **clothing_defects.sav.** File di dati ipotetici che prende in esame il processo di controllo di qualità di un'industria di abbigliamento. Per ciascun lotto prodotto nella fabbrica, gli ispettori prelevano un campione di abiti per contare il numero dei capi che non sono accettabili per la vendita.
- **coffee.sav.** Questo file di dati riguarda le immagini percepite di sei marchi di caffè ghiacciato⁶. Per ciascuno di 23 attributi di immagine del caffè freddo, le persone hanno selezionato tutti i marchi descritti dall'attributo. Le sei marche sono indicate dalle sigle AA, BB, CC, DD, EE e FF per tutelare la confidenzialità dei dati.
- **contacts.sav.** File di dati ipotetici che prende in esame l'elenco dei contatti di un gruppo di rappresentanti di vendita di computer aziendali. Ciascun contatto è classificato in base al reparto della società in cui lavora e dalle relative categorie aziendali. Il file riporta anche l'importo dell'ultima vendita effettuata, il tempo trascorso dall'ultima vendita e le dimensioni della società del contatto.
- **creditpromo.sav.** File di dati ipotetici che prende in esame l'impegno di un grande magazzino nel tentativo di valutare l'efficacia di una recente promozione con carta di credito. A tale scopo, sono stati selezionati 500 titolari di carta in modo casuale. Alla metà di questi è stato inviato un annuncio promozionale che comunica la riduzione del tasso d'interesse nel caso di acquisti effettuati entro i tre mesi successivi. All'altra metà è stato inviato un annuncio stagionale standard.
- **customer_dbase.sav.** File di dati ipotetici che prende in esame l'impegno di una società nel tentativo di utilizzare le informazioni contenute nel proprio database dei dati per creare offerte speciali per i clienti che più probabilmente risponderanno all'offerta. È stato selezionato in modo casuale un sottoinsieme della base dei clienti a cui è stata inviata l'offerta speciale e sono state registrate le risposte ricevute.
- **customer_information.sav.** File di dati ipotetici contenente le informazioni postali del cliente, ad esempio il nome e l'indirizzo.
- **customer_subset.sav.** Un sottoinsieme di 80 casi da *customer_dbase.sav*.
- **debate.sav.** File di dati ipotetici che prende in esame le risposte appaiate a un'indagine da parte dei partecipanti a un dibattito politico prima e dopo il dibattito. Ogni caso rappresenta un rispondente separato.

⁶ Kennedy, R., C. Riquier, e B. Affilato. 1996. Practical applications of correspondence analysis to categorical data in market research. *Journal of Targeting, Measurement, and Analysis for Marketing*, 5, 56-70.

- **debate_aggregate.sav.** File di dati ipotetici che aggrega le risposte contenute nel file *debate.sav*. Ciascun caso corrisponde a una classificazione incrociata della preferenza prima e dopo il dibattito.
- **demo.sav.** File di dati ipotetici che prende in esame un database di clienti che hanno fatto acquisti al fine di inviare offerte mensili tramite il metodo del direct mailing. Viene registrata la risposta dei clienti, sia che abbiano aderito all'offerta o meno, insieme a diverse informazioni demografiche.
- **demo_cs_1.sav.** File di dati ipotetici che prende in esame la prima fase che una società intraprende per compilare un database con informazioni ricavate dalle indagini. Ogni caso rappresenta una diversa città. Sono registrate anche le informazioni sulla regione, provincia, distretto e città.
- **demo_cs_2.sav.** File di dati ipotetici che prende in esame la seconda fase che una società intraprende per compilare un database con informazioni ricavate dalle indagini. Ogni caso rappresenta una diversa unità di abitazione, ricavata dalle città selezionate nella prima fase. Sono registrate anche le informazioni sulla regione, provincia, distretto, città, suddivisione e unità. Il file include inoltre informazioni sul campionamento ottenute dai primi due stadi del disegno.
- **demo_cs.sav.** File di dati ipotetici che contiene informazioni sulle indagini raccolte utilizzando un disegno di campionamento complesso. Ogni caso rappresenta una diversa unità di abitazione. Sono registrate diverse informazioni demografiche e sul campionamento.
- **diabetes_costs.sav.** Questo è un file di dati ipotetici che contiene informazioni gestite da una compagnia di assicurazioni su titolari di polizze che hanno il diabete. Ogni caso rappresenta un titolare di polizza differente.
- **dietstudy.sav.** Questo ipotetico file di dati contiene i risultati di uno studio della "dieta Stillman"⁷. Ogni caso corrisponde a un soggetto separato e registra i suoi pesi pregiati e post - dieta in chili e livelli di trigliceridi in mg/100 ml.
- **dmdata.sav.** Questo è un file di dati ipotetici che contiene informazioni demografiche e relative all'acquisto per una società di direct marketing. *dmdata2.sav* contiene informazioni per un sottoinsieme di contatti che hanno ricevuto un mailing di test e *dmdata3.sav* contiene informazioni sui restanti contatti che non hanno ricevuto il mailing di test.
- **dvdplayer.sav.** File di dati ipotetici che prende in esame lo sviluppo di un nuovo lettore DVD. Utilizzando un prototipo, il personale addetto al marketing ha raccolto dati sui gruppi di interesse. Ogni caso rappresenta un diverso utente che è stato sottoposto all'indagine e include informazioni demografiche personali dell'utente e sulle risposte che ha fornito riguardo al prototipo.
- **german_credit.sav.** Questo file di dati viene preso dal dataset "Credito tedesco" nel repository dei database Machine Learning⁸ presso l'Università della California, Irvine.
- **grocery_1month.sav.** Questo file di dati ipotetici corrisponde al file di dati *grocery_coupons.sav* con gli acquisti settimanali organizzati in modo che ogni caso corrisponda a un cliente separato. Alcune delle variabili che cambiano settimanalmente non vengono riportate nei risultati; l'importo speso registrato corrisponde ora alla somma degli importi spesi durante le quattro settimane dello studio.
- **grocery_coupons.sav.** File di dati ipotetici che contiene i dati di indagine raccolti da una catena di drogherie interessata alle abitudini di acquisto dei suoi clienti. Ciascun cliente viene seguito per quattro settimane e ciascun caso corrisponde a una settimana per cliente con informazioni sul luogo degli acquisti e i tipi di acquisti, incluso l'importo speso nelle drogherie durante la settimana.
- **guttman.sav.** Campana⁹ presentati un tavolo per illustrare possibili gruppi sociali. Guttman¹⁰ usato una porzione di questo tavolo, in cui cinque variabili che descrivono cose come l'interazione sociale, i sentimenti di appartenenza a un gruppo, la vicinanza fisica dei membri, e la formalità del rapporto sono stati incrociati con sette gruppi sociali teorici, tra cui la folla (ad esempio, le persone a una partita di calcio), le udienze (per esempio le persone a una lezione di teatro o di classe), pubbliche (ad esempio

⁷ Rickman, R., del contatto Mitchell, J. Dingman e J. E. Dalen. 1974. Changes in serum cholesterol during the Stillman Diet. *Journal of the American Medical Association*, 228:, 54-58.

⁸ Blake, C. L., e C. J. Merz. 1998. "UCI Repository of machine learning databases." Disponibile all'indirizzo <http://www.ics.uci.edu/~mllearn/MLRepository.html>.

⁹ E. H. Bell 1961. *Social foundations of human behavior: Introduction to the study of sociology*. New York: Harper & Row.

¹⁰ L. Guttman 1968. A general nonmetric technique for finding the smallest coordinate space for configurations of points. *Psychometrika*, 33, 469-506.

il pubblico o le audience televisive), le mob (come una folla ma con un'interazione molto più intensa), i gruppi primari (intimi), i gruppi secondari (volontari) e la comunità moderna (confederazione vaga risultante da una stretta prossimità fisica e da una necessità di servizi specializzati).

- **health_funding.sav.** File di dati ipotetici che contiene i dati sui fondi di assistenza sanitaria (importo per 100 persone), sui tassi di malattie (tasso per 10.000 persone) e sulle visite ai fornitori di assistenza sanitaria (tasso per 10.000 persone). Ogni caso rappresenta una diversa città.
- **hivassay.sav.** File di dati ipotetici che prende in esame l'impegno di un'industria farmaceutica nel tentativo di sviluppare un'analisi che riesca a rilevare in tempi brevi l'infezione da virus HIV. I risultati dell'analisi sono otto sfumature di colore rosso sempre più intenso; le sfumature più intense indicano la maggiore probabilità di infezione. Una prova di laboratorio è stata condotta su 2000 campioni di sangue. La metà di questi è risultata infetta al virus HIV, l'altra metà non è risultata infetta.
- **hourlywagedata.sav.** File di dati ipotetici che prende in esame la paga oraria degli infermieri occupati presso uffici e ospedali e in base ai diversi livelli di esperienza.
- **insurance_claims.sav.** Questo è un file di dati ipotetici che prende in esame una compagnia di assicurazioni che desidera creare un modello per segnalare richieste di risarcimento sospette e potenzialmente fraudolente. Ogni caso rappresenta una richiesta di risarcimento separata.
- **insure.sav.** File di dati ipotetici che prende in esame una compagnia di assicurazioni impegnata nello studio dei fattori di rischio, che indicano l'eventualità che un cliente presenti una domanda di indennizzo in un contratto assicurativo sulla vita della durata di dieci anni. Ogni caso nel file di dati rappresenta una coppia di contratti. In un contratto sono contenute informazioni su una richiesta di risarcimento, l'altro sull'età e sul sesso.
- **judges.sav.** File di dati ipotetici che prende in esame il punteggio assegnato, da giurie qualificate (più un appassionato) a 300 prestazioni sportive. Ciascuna riga rappresenta una diversa prestazione; i giudici hanno esaminato la stessa prestazione.
- **kinship_dat.sav.** Rosenberg e Kim¹¹ set per analizzare 15 termini di kinship (zia, fratello, cugino, figlia, padre, nipote, nonno, nonna, nipote, madre, nipote, nipote, sorella, figlio, zio). Hanno richiesto a quattro gruppi di studenti universitari (due composti da femmine e due da maschi) di ordinare questi termini in base alla similarità. A due gruppi (uno femminile e uno maschile) è stato richiesto di effettuare l'ordinamento due volte, con il secondo ordinamento basato su un criterio diverso rispetto al primo. Pertanto, è stato ottenuto un totale di sei "origini". Ogni origine corrisponde ad una matrice di prossimità 15 x 15, le cui celle sono uguali al numero di persone in un'origine meno il numero di volte in cui gli oggetti sono stati partizionati insieme in quell'origine.
- **kinship_ini.sav.** Questo file di dati contiene la configurazione iniziale di una soluzione a tre dimensioni per *kinship_dat.sav*.
- **kinship_var.sav.** Questo file di dati contiene variabili indipendenti relative a sesso, generazione e grado di separazione che possono essere utilizzate per interpretare le dimensioni di una soluzione per *kinship_dat.sav*. In modo specifico, tali variabili possono essere utilizzate per limitare lo spazio della soluzione a una combinazione lineare di tali variabili.
- **marketvalues.sav.** Questo file di dati riguarda le vendite a domicilio in un nuovo sviluppo abitativo in Algonquin, Ill., durante gli anni dal 1999-2000. Queste vendite sono una questione di record pubblico.
- **nhis2000_subset.sav.** Il National Health Interview Survey (NHIS) è un'indagine di grandi dimensioni condotta sulla popolazione civile americana. Le interviste vengono realizzate di persona e si basano su un campione rappresentativo di famiglie a livello nazionale. Per ogni membro di una famiglia vengono raccolte osservazioni e informazioni di carattere demografico relative allo stato di salute. Questo file di dati contiene un sottoinsieme delle informazioni ottenute dall'indagine del 2000. National Center for Health Statistics. National Health Interview Survey, 2000. File di dati e documentazione di utilizzo pubblico. ftp://ftp.cdc.gov/pub/Health_Statistics/NCHS/Datasets/NHIS/2000/. Accesso 2003.
- **ozone.sav.** I dati includono 330 osservazioni basate su sei variabili meteorologiche per quantificare la concentrazione dell'ozono dalle variabili rimanenti. Ricercatori precedenti^{12,13}, tra le altre trovate non linearità tra queste variabili, che ostacolano approcci di regressione standard.

¹¹ Rosenberg, S., e M. P. Kim. 1975. The method of sorting as a data-gathering procedure in multivariate research. *Multivariate Behavioral Research*, 10, 489-502.

- **pain_medication.sav.** File di dati ipotetici che contiene i risultati di un test clinico per stabilire la cura antinfiammatoria per il trattamento del dolore generato dall'artrite cronica. Di particolare interesse, il test ha evidenziato il tempo che impiega il farmaco ad avere effetto e il confronto con altri farmaci esistenti.
- **patient_los.sav.** File di dati ipotetici che contiene informazioni sul trattamento dei pazienti ricoverati per sospetto di infarto del miocardio. Ogni caso corrisponde a un diverso paziente e contiene diverse variabili correlate alla degenza nell'ospedale.
- **patlos_sample.sav.** File di dati ipotetici che contiene informazioni sul trattamento di un campione di pazienti curato con trombolitici durante la degenza per infarto del miocardio. Ogni caso corrisponde a un diverso paziente e contiene diverse variabili correlate alla degenza nell'ospedale.
- **poll_cs.sav.** File di dati ipotetici che prende in esame i sondaggi per stabilire il livello di sostegno pubblico nei confronti di un disegno di legge prima che diventi una legge vera e propria. I casi corrispondono ai votanti registrati. Ciascun caso riporta informazioni sulla contea, sul comune e sul quartiere in cui vive il votante.
- **poll_cs_sample.sav.** File di dati ipotetici che contiene un campione dei votanti elencati nel file *poll_cs.sav*. Il campione è stato prelevato in base al disegno specificato nel file di piano *poll.csplan* e questo file di dati contiene le probabilità di inclusione e i pesi del campione. Tuttavia, notare che poiché fa uso del metodo PPS (probability-proportional-to-size, probabilità proporzionale alla dimensione), esiste anche un file contenente le probabilità di selezione congiunte (*poll_jointprob.sav*). Le ulteriori variabili corrispondenti ai dati demografici dei votanti e alla loro opinione sul disegno di legge, sono state raccolte e aggiunte al file di dati dopo aver acquisito il campione.
- **property_assess.sav.** File di dati ipotetici che prende in esame l'impegno di un perito di una contea nel tentativo di mantenere gli accertamenti sui valori delle proprietà aggiornati in base alle risorse limitate. I casi rappresentano le proprietà vendute nella contea nello scorso anno. Ogni caso nel file di dati contiene informazioni sul comune in cui si trova la proprietà, il perito che per ultimo ha visitato la proprietà, il tempo trascorso dall'accertamento, la valutazione fatta in tale momento e il valore di vendita della proprietà.
- **property_assess_cs.sav.** File di dati ipotetici che prende in esame l'impegno di un perito di uno stato nel tentativo di mantenere aggiornati gli accertamenti sui valori delle proprietà in base alle risorse limitate. I casi corrispondono alle proprietà nello stato. Ogni caso nel file di dati include informazioni sulla contea, il comune e il quartiere in cui risiede la proprietà, la data dell'ultimo accertamento e la valutazione fatta in tale data.
- **property_assess_cs_sample.sav.** File di dati ipotetici che contiene un campione delle proprietà elencate nel file *property_assess_cs.sav*. Il campione è stato prelevato in base al disegno specificato nel file di piano *property_assess.csplan* e questo file di dati contiene le probabilità di inclusione e i pesi del campione. L'ulteriore variabile *Valore corrente* è stata raccolta e aggiunta al file di dati dopo aver acquisito il campione.
- **recidivism.sav.** File di dati ipotetici che prende in esame l'impegno delle Forze dell'Ordine nel tentativo di valutare il tasso di recidività nella propria area di giurisdizione. Ogni caso corrisponde a un precedente trasgressore e include le informazioni demografiche, alcuni dettagli sul primo crimine, il tempo trascorso fino al secondo arresto e se tale arresto è avvenuto entro due anni dal primo.
- **recidivism_cs_sample.sav.** File di dati ipotetici che prende in esame l'impegno delle Forze dell'Ordine nel tentativo di valutare il tasso di recidività nella propria area di giurisdizione. Ogni caso corrisponde a un precedente trasgressore, rilasciato dopo il primo arresto durante il mese di giugno 2003, e include le informazioni demografiche, alcuni dettagli sul primo crimine e i dati del secondo arresto, se si è verificato entro la fine di giugno 2006. I trasgressori sono stati selezionati da dipartimenti campione in base al piano di campionamento specificato in *recidivism_cs.csplan*; poiché fa uso del metodo PPS (probability-proportional-to-size, probabilità proporzionale alla dimensione), esiste anche un file contenente le probabilità di selezione congiunte (*recidivism_cs_jointprob.sav*).

¹² Breiman, L., e J. H. Friedman. 1985. Estimating optimal transformations for multiple regression and correlation. *Journal of the American Statistical Association*, 80, 580-598.

¹³ Hastie, T., e R. Tibshirani. 1990. *Generalized additive models*. London: Chapman and Hall.

- **rfm_transactions.sav.** File di dati ipotetici contenente i dati della transazione di acquisto, inclusa la data di acquisto, gli articoli acquistati e il valore monetario di ciascuna transazione.
- **salesperformance.sav.** File di dati ipotetici che prende in esame la valutazione di due nuovi corsi di formazione alle vendite. Sessanta dipendenti, divisi in tre gruppi, ricevono tutti la formazione standard. In più, al gruppo 2 viene assegnato un corso di formazione tecnica e al gruppo 3 un'esercitazione pratica. Al termine del corso di formazione, ciascun dipendente viene sottoposto a un esame e il punteggio conseguito viene registrato. Ciascun caso nel file di dati rappresenta un diverso partecipante. Il file di dati include il gruppo a cui è assegnato il partecipante e il punteggio conseguito all'esame finale.
- **satisf.sav.** File di dati ipotetici che prende in esame un'indagine sulla soddisfazione dei clienti condotta da una società di vendita al dettaglio presso 4 negozi. Sono stati intervistati 582 clienti e ciascun caso rappresenta le risposte ottenute da un singolo cliente.
- **screws.sav.** Questo file di dati contiene informazioni sulle caratteristiche delle viti, bulloni, dadi e tacchi¹⁴.
- **shampoo_ph.sav.** File di dati ipotetici che prende in esame il processo di controllo di qualità di un'industria di prodotti per capelli. A intervalli di tempo regolari, vengono misurati sei diversi lotti prodotti e ne viene registrato il relativo pH. L'intervallo di obiettivo è 4.5–5.5.
- **ships.sav.** Un dataset presentato e analizzato altrove¹⁵ che riguarda i danni alle navi da carico causate dalle onde. I conteggi degli incidenti possono essere presentati con un tasso di Poisson in base al tipo di nave, al periodo di costruzione e al periodo di servizio. I mesi di servizio aggregati di ciascuna cella della tabella generata dalla classificazione incrociata dei fattori fornisce i valori di esposizione al rischio.
- **site.sav.** File di dati ipotetici che prende in esame l'impegno di una società nella scelta di nuovi siti in cui espandere la propria presenza. La società ha incaricato due consulenti separati che, oltre a valutare i siti e presentare un report completo, devono classificarli come potenzialmente "molto adatti", "adatti" o "poco adatti".
- **smokers.sav.** Questo file di dati è astratto dal National Household Survey of Drug Abuse ed è un campione di probabilità delle famiglie americane. (<http://dx.doi.org/10.3886/ICPSR02934>) Così, il primo passo di un'analisi di questo file di dati dovrebbe essere quello di mettere in peso i dati per riflettere le tendenze della popolazione.
- **stocks.sav** Questo file di dati ipotetici contiene i prezzi delle azioni e il volume per un anno.
- **stroke_clean.sav.** Questo file di dati ipotetici contiene lo stato di un database medico dopo averne eseguito la pulizia utilizzando le procedure nell'opzione Data Preparation.
- **stroke_invalid.sav.** Questo file di dati ipotetici contiene riporta lo stato iniziale di un database medico e contiene numerosi errori di immissione dati.
- **stroke_survival.** Questo file di dati ipotetici riguarda i tempi di sopravvivenza per i pazienti che, dopo avere completato un programma riabilitativo in seguito a un ictus postischemico, affrontano alcune sfide. Dopo l'attacco, viene annotata l'occorrenza dell'infarto miocardico, dell'ictus ischemico o emorragico e viene registrata l'ora dell'evento. Questo campione viene troncato a sinistra perché include solo i pazienti che sono sopravvissuti fino al termine del programma riabilitativo post-ictus.
- **stroke_valid.sav.** File di dati ipotetici che riporta lo stato di un database medico dopo il controllo dei valori eseguito con la procedura Convalida i dati. Il database contiene comunque casi potenzialmente anomali.
- **survey_sample.sav.** File di dati che contiene i dati dell'indagine, compresi i dati demografici e varie misure dell'atteggiamento. Si basa su un sottoinsieme di variabili tratte dal 1998 NORC General Social Survey, benché alcuni valori dei dati siano stati modificati e siano state aggiunte variabili dummy a scopo dimostrativo.
- **tcm_kpi.sav.** Questo è un file di dati ipotetici che contiene valori di indicatori di prestazioni chiave settimanali per un business. Contiene anche i dati settimanali per un numero di metriche controllabili durante lo stesso periodo di tempo.

¹⁴ J. A. Hartigan 1975. *Clustering algorithms*. New York: John Wiley and Sons.

¹⁵ McCullagh, P., e J. A. Nelder. 1989. *Modelli lineari generalizzati*, 2a ed. Londra: Chapman & Hall.

- **tcm_kpi_upd.sav.** Questo file di dati è identico al file *tcm_kpi.sav*, ma contiene i dati per quattro settimane extra.
- **telco.sav.** File di dati ipotetici che prende in esame l'impegno di un'azienda di telecomunicazioni nel tentativo di ridurre il tasso di abbandono, ovvero l'abbandono dei propri clienti. Ciascun caso rappresenta un cliente separato e riporta diverse informazioni demografiche e sull'uso del servizio.
- **telco_extra.sav.** Questo file di dati è simile al file *telco.sav*, ma le variabili "tenure" e spesa del cliente trasformata tramite logaritmo sono state sostituite dalle variabili di spesa del cliente trasformata tramite logaritmo standardizzate.
- **telco_missing.sav.** Questo file di dati è un sottoinsieme del file di dati *telco.sav*, ma alcuni dei valori dei dati demografici sono stati sostituiti con valori mancanti.
- **testmarket.sav.** File di dati ipotetici che prende in esame i piani di una catena di fast food per aggiungere un nuovo prodotto al proprio menu. Sono previste tre campagne promozionali del nuovo prodotto. Il prodotto viene introdotto in diversi mercati selezionati in modo casuale. Per ogni sede viene utilizzata una promozione differente registrando le vendite settimanali della nuova voce per le prime quattro settimane. Ogni caso rappresenta un luogo e una settimana diversi.
- **testmarket_1month.sav.** Questo file di dati ipotetici corrisponde al file *testmarket.sav* con le vendite settimanali organizzate in modo che ogni caso corrisponda a un luogo separato. Alcune delle variabili che cambiano settimanalmente non vengono riportate nei risultati; le vendite registrate corrispondono ora alla somma delle vendite conseguite durante le quattro settimane dello studio.
- **tree_car.sav.** File di dati ipotetici che contiene dati demografici e sul prezzo di acquisto dei veicoli.
- **tree_credit.sav.** File di dati ipotetici che contiene dati demografici e sulla cronologia dei mutui di una banca.
- **tree_missing_data.sav.** File di dati ipotetici che contiene dati demografici e sulla cronologia dei mutui di una banca con un numero elevato di valori mancanti.
- **tree_score_car.sav.** File di dati ipotetici che contiene dati demografici e sul prezzo di acquisto dei veicoli.
- **tree_textdata.sav.** File di dati semplice con due variabili destinato principalmente per mostrare lo stato predefinito delle variabili prima dell'assegnazione dei livelli di misurazione e delle etichette valore.
- **tv-survey.sav.** File di dati ipotetici che prende in esame un'indagine condotta da una emittente televisiva che deve stabilire se estendere la durata di un programma di successo. A un campione di 906 rispondenti è stato chiesto se preferisce guardare il programma con diverse condizioni. Ciascuna riga rappresenta un diverso rispondente e ciascuna colonna una diversa condizione.
- **ulcer_recurrence.sav.** Questo file contiene informazioni parziali su uno studio svolto per mettere a confronto l'efficacia di due terapie preventive per la recidiva delle ulcere. Fornisce un buon esempio di intervallo - dati censurati ed è stato presentato e analizzato altrove¹⁶.
- **ulcer_recurrence_recoded.sav.** In questo file sono contenute le informazioni del file *ulcer_recurrence.sav* riorganizzate per consentire di presentare la probabilità degli eventi per ciascun intervallo dello studio, anziché solo al termine. È stato presentato e analizzato altrove¹⁷.
- **verd1985.sav.** Questo file di dati riguarda un sondaggio¹⁸. Sono state registrate le risposte di 15 soggetti a 8 variabili. Le variabili di interesse sono suddivise in tre insiemi. L'insieme 1 include *età* e *statociv*, l'insieme 2 include *andom* e *giornale* e l'insieme 3 include *musica* e *vicinato*. *Andom* viene scalata come nominale multipla ed *età* come ordinale; tutte le altre variabili vengono scalate come nominali singole.
- **virus.sav.** File di dati ipotetici che prende in esame l'impegno di un ISP (Internet Service Provider) nel tentativo di determinare gli effetti che un virus può generare nelle sue reti. Si è tenuta traccia della

¹⁶ D. Collett 2003. *Modellazione dati di sopravvivenza nella ricerca medica*, 2 ed. Boca Raton: Chapman & Hall / CRC.

¹⁷ D. Collett 2003. *Modellazione dati di sopravvivenza nella ricerca medica*, 2 ed. Boca Raton: Chapman & Hall / CRC.

¹⁸ R. Verdegaal 1985. *Meer sets analyse voor kwalitatieve gegevens (in olandese)*. Leiden: Department of Data Theory, University of Leiden.

percentuale (approssimativa) di traffico e-mail infettato da virus sulla rete in un lasso di tempo, dal momento dell'individuazione fino alla soppressione della minaccia.

- **wheeze_steubenville.sav.** Si tratta di un sottoinsieme da uno studio longitudinale degli effetti sulla salute dell'inquinamento dell'aria sui bambini¹⁹. I dati contengono misure binarie ripetute dello status di respiro sismico per i bambini di Steubenville, Ohio, a 7 anni, 8, 9 e 10 anni, insieme a una registrazione fissa di se la madre fosse o meno fumatrice durante il primo anno di studio.
- **workprog.sav.** File di dati ipotetici che prende in esame un programma di lavoro governativo il cui obiettivo è fornire attività più adatte alle persone diversamente abili. È stato seguito un campione di potenziali partecipanti al programma, alcuni dei quali sono stati selezionati in modo casuale e altri no. Ogni caso rappresenta un diverso partecipante al programma.
- **worldsales.sav** Questo file di dati ipotetici contiene i ricavi delle vendite per continente e prodotto.

¹⁹ Ware, J. H., D. W. Dockery, A. Spiro III, F. E. Speizer e B. G. Ferris Junior. 1984. Passive smoking, gas cooking, and respiratory health of children living in six cities. *American Review of Respiratory Diseases*, 129, 366-374.

Informazioni particolari

Queste informazioni sono state sviluppate per prodotti e servizi offerti negli Stati Uniti. Questo materiale potrebbe essere disponibile da IBM in altre lingue. Tuttavia, può essere necessario possedere una copia del prodotto o versione del prodotto in tale lingua per potervi accedere.

IBM può non offrire i prodotti, i servizi o le funzioni presentati in questo documento in altri paesi. Consultare il proprio rappresentante locale IBM per informazioni sui prodotti ed i servizi attualmente disponibili nella propria zona. Qualsiasi riferimento ad un prodotto, programma o servizio IBM non implica o intende dichiarare che solo quel prodotto, programma o servizio IBM può essere utilizzato. Qualsiasi prodotto funzionalmente equivalente al prodotto, programma o servizio che non violi alcun diritto di proprietà intellettuale IBM può essere utilizzato. È comunque responsabilità dell'utente valutare e verificare la possibilità di utilizzare altri prodotti, programmi o servizi non IBM.

IBM può avere applicazioni di brevetti o brevetti in corso relativi all'argomento descritto in questo documento. La fornitura del presente documento non concede alcuna licenza a tali brevetti. È possibile inviare per iscritto richieste di licenze a:

*IBM Director of Licensing
IBM Corporation
North Castle Drive, MD-NC119
Armonk, NY 10504-1785
it*

Per richieste di licenze relative ad informazioni double-byte (DBCS), contattare il Dipartimento di Proprietà Intellettuale IBM nel proprio paese o inviare richieste per iscritto a:

*Intellectual Property Licensing
Legal and Intellectual Property Law
IBM Japan Ltd.
19-21, Nihonbashi-Hakozakicho, Chuo-ku
Tokyo 103-8510, Japan*

IBM (INTERNATIONAL BUSINESS MACHINES CORPORATION) FORNISCE LA PRESENTE PUBBLICAZIONE "NELLO STATO IN CUI SI TROVA" SENZA GARANZIE DI ALCUN TIPO, ESPRESSE O IMPLICITE, IVI INCLUSE, A TITOLO DI ESEMPIO, GARANZIE IMPLICITE DI NON VIOLAZIONE, DI COMMERCIALIZZABILITÀ E DI IDONEITÀ PER UNO SCOPO PARTICOLARE. Alcune giurisdizioni non consentono la rinuncia ad alcune garanzie espresse o implicite in determinate transazioni, pertanto, la presente dichiarazione può non essere applicabile.

Queste informazioni potrebbero contenere imprecisioni tecniche o errori tipografici. Le modifiche alle presenti informazioni vengono effettuate periodicamente; tali modifiche saranno incorporate nelle nuove pubblicazioni della pubblicazione. IBM si riserva il diritto di apportare miglioramenti e/o modifiche al prodotto o al programma descritto nel manuale in qualsiasi momento e senza preavviso.

I riferimenti in queste informazioni a siti Web non IBM vengono forniti solo per comodità e non implicano in alcun modo l'approvazione di tali siti web. I materiali presenti in tali siti web non sono parte dei materiali per questo prodotto IBM e l'utilizzo di tali siti è a proprio rischio.

IBM può utilizzare o distribuire qualsiasi informazione fornita in qualsiasi modo ritenga appropriato senza incorrere in alcun obbligo verso l'utente.

Coloro che detengano la licenza su questo programma e desiderano avere informazioni su di esso allo scopo di consentire: (i) uno scambio di informazioni tra programmi indipendenti ed altri (compreso questo) e (ii) l'utilizzo reciproco di tali informazioni, dovrebbe rivolgersi a:

*IBM Director of Licensing
IBM Corporation
North Castle Drive, MD-NC119*

Tali informazioni possono essere disponibili, in base ad appropriate clausole e condizioni, includendo in alcuni casi, il pagamento di una tassa.

Il programma concesso in licenza descritto nel presente documento e tutto il materiale concesso in licenza disponibile sono forniti da IBM in base alle clausole dell'Accordo per Clienti IBM (IBM Customer Agreement), dell'IBM IPLA (IBM International Program License Agreement) o qualsiasi altro accordo equivalente tra le parti.

Gli esempi client e dati delle prestazioni citati sono presentati esclusivamente a scopo illustrativo. Gli effettivi risultati delle prestazioni possono variare in base alle configurazioni e alle condizioni operative specifiche.

Le informazioni relative a prodotti non IBM sono ottenute dai fornitori di quei prodotti, dagli annunci pubblicati o da altre fonti disponibili al pubblico. IBM non ha testato quei prodotti e non può confermarne la precisione della prestazione, la compatibilità o qualsiasi altro reclamo relativo ai prodotti non IBM. Le domande sulle funzionalità dei prodotti non IBM devono essere indirizzate ai fornitori di tali prodotti.

Le dichiarazioni relative all'orientamento o alle intenzioni future di IBM sono soggette a modifica o a ritiro senza preavviso e rappresentano solo mete e obiettivi.

Questa pubblicazione contiene esempi di dati e prospetti utilizzati quotidianamente nelle operazioni aziendali. Pertanto, per maggiore completezza, gli esempi includono nomi di persone, società, marchi e prodotti. Tutti i nomi contenuti nel manuale sono fittizi e ogni riferimento a persone o aziende reali è puramente casuale.

LICENZA DI COPYRIGHT:

Queste informazioni contengono programmi applicativi di esempio in linguaggio sorgente, che illustrano tecniche di programmazione su varie piattaforme operative. È possibile copiare, modificare e distribuire questi programmi di esempio sotto qualsiasi forma senza alcun pagamento a IBM, allo scopo di sviluppare, utilizzare, commercializzare o distribuire i programmi applicativi in conformità alle API (application programming interface) a seconda della piattaforma operativa per cui i programmi di esempio sono stati scritti. Questi esempi non sono stati testati approfonditamente tenendo conto di tutte le condizioni possibili. IBM, quindi, non può garantire o sottintendere l'affidabilità, l'utilità o il funzionamento di questi programmi. I programmi di esempio sono forniti "NELLO STATO IN CUI SI TROVANO", senza alcun tipo di garanzia. IBM non intende essere responsabile per alcun danno derivante dall'uso dei programmi di esempio.

Ogni copia o qualsiasi parte di questi programmi di esempio o qualsiasi lavoro derivato, devono contenere le seguenti informazioni relative alle leggi sul diritto d'autore:

© Copyright IBM Corp. 2021. Le porzioni di questo codice derivano da IBM Corp. Programmi di esempio.

© Copyright IBM Corp. 1989 - 2021. Tutti i diritti riservati.

Marchi

IBM, il logo IBM e ibm.com sono marchi o marchi registrati di International Business Machines Corp., registrati in molte giurisdizioni in tutto il mondo. Altri nomi di prodotti e servizi possono essere marchi di IBM o di altre società. Un elenco corrente dei marchi IBM è disponibile sul web in "Copyright and trademark information" all'indirizzo www.ibm.com/legal/copytrade.shtml.

Adobe, il logo Adobe, PostScript e il logo PostScript sono marchi o marchi registrati di Adobe Systems Incorporated negli Stati Uniti e/o in altri paesi.

Intel, il logo Intel, Intel Inside, il logo Intel Inside, Intel Centrino, il logo Intel Centrino, Celeron, Intel Xeon, Intel SpeedStep, Itanium e Pentium sono marchi o marchi registrati di Intel Corporation o delle sue consociate negli Stati Uniti e in altri paesi.

Linux è un marchio registrato di Linus Torvalds negli Stati Uniti e/o negli altri paesi.

Microsoft, Windows, Windows NT e il logo Windows sono marchi di Microsoft Corporation negli Stati Uniti e/o negli altri paesi.

UNIX è un marchio della The Open Group negli Stati Uniti e/o negli altri paesi.

Java e tutti i marchi e i logo basati su Java sono marchi o marchi registrati di Oracle e/o associate.

Indice analitico

C

- categorie calcolate
 - da subtotali [37](#)
 - formati di visualizzazione [14](#)
 - nascondere categorie in espressione [36](#)
 - Tabelle personalizzate [13](#), [35](#)
- categorie post - elaborazioni
 - Tabelle personalizzate [13](#), [35](#)
- celle vuote
 - valore visualizzato nelle tabelle personalizzate [15](#), [75](#)
- chi-quadrato
 - Tabelle personalizzate [53](#)
- colonna significa statistiche
 - tabelle personalizzate [55](#)
- compressione delle categorie
 - Tabelle personalizzate [34](#)
- conteggio
 - vs. valido N [48](#)
- controllo numero di decimali visualizzati [22](#)

D

- data
 - inclusa la data corrente nelle tabelle personalizzate [16](#)
- deviazione standard
 - Tabelle personalizzate [9](#)
- didascalie
 - Tabelle personalizzate [16](#)

E

- elaborazione file suddiviso
 - tabelle personalizzate [4](#)
- Eliminazione di categorie
 - Tabelle personalizzate [11](#), [23](#)
- esclusione di categorie
 - Tabelle personalizzate [11](#), [23](#)
- etichette
 - modifica testo etichetta per le statistiche di riepilogo [74](#)
- Etichette di angolo
 - Tabelle personalizzate [16](#)
- etichette di variabile
 - sopprimendo la visualizzazione nelle tabelle personalizzate [5](#)

F

- file di esempio
 - ubicazione [76](#)
- formati di visualizzazione
 - statistiche di riepilogo in tabelle personalizzate [10](#), [72](#)

I

- insiemi a risposta multipla

- insiemi a risposta multipla (*Continua*)
 - percentuali [8](#)
 - risposte duplicate in più insiemi di categoria [15](#)
 - test di significato [63](#), [68](#), [69](#)
- intervalli di confidenza [52](#)

L

- larghezza colonna
 - controllo nelle tabelle personalizzate [15](#), [74](#)
- livello di misurazione
 - modifica nelle tabelle personalizzate [1](#)

M

- massimo
 - Tabelle personalizzate [9](#)
- media
 - Tabelle personalizzate [9](#)
- mediana
 - Tabelle personalizzate [9](#)
- minimo
 - Tabelle personalizzate [9](#)
- modo
 - Tabelle personalizzate [9](#)

N

- nascondere le etichette statistiche nelle tabelle personalizzate [19](#)
- Numero di casi validi
 - Tabelle personalizzate [9](#)
- Numero totale di casi [71](#)

O

- omettere categorie
 - Tabelle personalizzate [23](#)
- ora
 - incluso il tempo corrente nelle tabelle personalizzate [16](#)
- ordinamento di categorie
 - Tabelle personalizzate [23](#)

P

- percentuali
 - insiemi a risposta multipla [8](#)
 - nelle tabelle personalizzate [7](#), [8](#), [20](#), [21](#)
 - valori mancanti [71](#)
- posizioni decimali
 - controllo del numero di decimali visualizzati nelle tabelle personalizzate [6](#), [22](#), [72](#)

R

- riassunti raggruppati

riassunti raggruppati (*Continua*)
variabili di scala [50](#)
riordino categorie
Tabelle personalizzate [11](#)

S

scala
Tabelle personalizzate [9](#)
somma
Tabelle personalizzate [9](#)
stampa tabelle con strati [30](#)
Statistiche
statistiche di riepilogo totale personalizzate [45](#)
statistiche riassuntive [42](#)
tabelle in sovrapposizione [45](#)
statistiche delle proporzioni di colonna
tabelle personalizzate [58](#)
statistiche di riepilogo diverse per diverse variabili
tabelle in sovrapposizione [49](#)
statistiche di riepilogo totale personalizzate [45](#)
Statistiche di test
Tabelle personalizzate [17, 53](#)
statistiche riassuntive
dimensione origine [43](#)
formato di visualizzazione [72](#)
modifica testo etichetta [74](#)
riepiloghi differenti per diverse variabili in tabelle
impilate [49](#)
statistiche di riepilogo totale personalizzate [45](#)
tabelle in sovrapposizione [45](#)
variabile di origine [43](#)

T

tabella delle frequenze
Tabelle personalizzate [15, 39](#)
tabelle
Tabelle personalizzate [1](#)
tabelle di comperimetro [15, 39](#)
tabelle di frequenza media [10, 45](#)
tabelle personalizzate
elaborazione file suddiviso [4](#)
Tabelle personalizzate
categorie calcolate [11, 13, 35](#)
categorie post - elaborazioni [13, 35](#)
celle vuote [15](#)
come costruire una tabella [3](#)
compressione delle categorie [34](#)
controllo numero di decimali visualizzati [6](#)
didascalie [16](#)
dimensione di origine statistiche [21](#)
esclusione di categorie [11, 23](#)
Esclusione valori mancanti per riepiloghi di scala [15](#)
Etichette di angolo [16](#)
etichette di valore per variabili categoriali [1](#)
formati di visualizzazione [6](#)
Formati di visualizzazione delle statistiche riassuntive
[10](#)
insiemi a categorie multiple [15](#)
insiemi a risposta multipla [1, 63](#)
larghezza colonna [15](#)
modifica del livello di misurazione [1](#)

Tabelle personalizzate (*Continua*)
modifica delle etichette per statistiche di riepilogo [20](#)
modifica dimensione statistiche di riepilogo [10](#)
mostrare e nascondere nomi e etichette variabili [5](#)
nascondere categorie sottototali [34](#)
nascondere etichette statistiche [19](#)
ordinamento di categorie [23](#)
oscillazione delle variabili della riga e della colonna [28](#)
percentuali [7, 8, 20, 21](#)
percentuali per più set di risposta [8](#)
riga rispetto alle percentuali di colonna [20](#)
riordino categorie [11](#)
stampa tabelle stratificate [30](#)
Statistiche di test [17, 53](#)
statistiche riassuntive [7–9](#)
tabella delle frequenze [15, 39](#)
tabelle di comperimetro [15, 39](#)
tabelle di frequenza media [10](#)
tabelle di variabili con categorie condivise [15, 39](#)
tabelle semplici per variabili categoriali [19](#)
tavola di contingenza [21](#)
test di significatività e risposta multipla [63](#)
titoli [16](#)
totali [11, 20, 30](#)
totali in tabelle con categorie escluse [23](#)
totali marginali [22](#)
totali parziali [11, 30](#)
totali personalizzati [10](#)
Variabili categoriali [1](#)
variabili di nesting [25, 27](#)
variabili di scala [1](#)
variabili di stack [24, 25](#)
Variabili di strato [28–30](#)
variabili strato di nidificazione [30](#)
Vista compatta [27](#)
tavola di contingenza
Tabelle personalizzate [21](#)
Test di significatività
insiemi a risposta multipla [68, 69](#)
Tabelle personalizzate [17](#)
titoli
Tabelle personalizzate [16](#)
totali
categorie escluse [31](#)
posizione di visualizzazione [31](#)
strati [33](#)
tabelle nidificate [32](#)
Tabelle personalizzate [11, 20, 30](#)
totali di gruppo [32](#)
totali marginali per tabelle personalizzate [22](#)
totali del sottogruppo [32](#)
totali di gruppo [32](#)
totali parziali
nascondere categorie sottototali [34](#)
Tabelle personalizzate [11, 30](#)

V

valori
visualizzazione di etichette di categoria e valori [46](#)
valori e etichette di valore [46](#)
valori mancanti
effetto sui calcoli percentuali [71](#)
incluso nelle tabelle personalizzate [71](#)

- valori mancanti definiti dall'utente [69](#)
- valori mancanti di sistema [69](#)
- variabile di origine statistiche di riepilogo
 - variabili di scala [51](#)
- variabili di nesting
 - Tabelle personalizzate [25](#), [27](#)
 - variabili di scala [51](#)
- variabili di scala
 - annidamento [51](#)
 - esecuzione stack [48](#)
 - riassunti raggruppati [50](#)
 - riepiloghi raggruppati per riga e colonne categoriali [51](#)
 - statistiche di riepilogo multiple [48](#)
 - statistiche riassuntive [48](#)
- variabili di stack
 - statistiche di riepilogo diverse per diverse variabili [49](#)
 - Tabelle personalizzate [24](#), [25](#)
 - variabili di origine statistiche di riepilogo multipla [45](#)
 - variabili di scala [48](#)
 - variabili layer di stack [29](#)
- Variabili di strato
 - stampa tabelle stratificate [30](#)
 - Tabelle personalizzate [28–30](#)
 - variabili layer di stack [29](#)
 - variabili strato di nidificazione [30](#)
- varianza
 - Tabelle personalizzate [9](#)
- visualizzazione valori di categoria [46](#)

