

UNIVERSITÀ DEGLI STUDI DI NAPOLI FEDERICO II



Dipartimento di Ingegneria Civile, Edile e Ambientale

Dottorato di Ricerca in
Ingegneria dei Sistemi Civili XXXIV ciclo

Dottorando
Buonocore Ciro

Titolo della tesi di ricerca
***Metodi innovativi di stima e correzione della matrice origine-
destinazione con conteggi di traffico***

Coordinatore di dottorato:

Prof. Ing. Andrea Papola

Tutor:

Prof. Ing. Andrea Papola

Sommario

1.	Introduzione	7
1.1.	Background	7
1.2.	Outline.....	9
2.	Stima della domanda: generalità	11
2.1.	Studio di un sistema di trasporto.....	11
2.2.	Stima della domanda di mobilità: stima diretta	16
2.3.	Stima della domanda di mobilità: modelli analitici	19
2.3.1.	Modelli basati sull'utilità aleatoria	21
2.3.2.	Modelli descrittivi: modello gravitazionale.....	24
2.4.	Correzione dei flussi <i>od</i> mediante conteggi di traffico.....	25
2.4.1.	Approcci per ridurre lo sbilancio equazioni incognite	28
2.4.2.	Approcci per garantire la coverage.....	29
2.4.3.	Stimatori per la correzione della domanda.....	31
3.	Correzione della matrice di domanda mediante clustering sequenziale di flussi <i>od</i> ..	35
3.1.	Descrizione della metodologia.....	36
3.2.	Descrizione del caso studio.....	41
3.3.	Risultati e conclusioni.....	44
4.	Correzione della matrice di domanda mediante clustering simultaneo di flussi <i>od</i> ..	56
4.1.	Descrizione metodologia	56
4.2.	Descrizione caso studio.....	60
4.2.1.	Setting 1	60

4.2.2. Setting 2	62
4.3. Risultati e conclusioni	64
5. Correzione della matrice di domanda mediante modello gravitazionale q-generalizzato e autocorrelazioni spaziali basate sulle adiacenze	71
5.1. Descrizione metodologia	73
5.2. Descrizione caso studio.....	77
5.3. Risultati e conclusioni.....	80
6. Conclusioni e sviluppi futuri	108
7. Bibliografia.....	113

Indice figure

Fig. 1.1	Zonizzazione	12
Fig. 1.2	Matrice od	12
Fig. 2.1	Classificazione dei modelli di domanda	21
Fig. 3.1	Esempio di clusterizzazione su rete bidirezionale con 2 sezioni di conteggio ..	38
Fig. 3.2	Schema della metodologia	40
Fig. 3.3	Rete test.....	41
Fig. 3.4	Rete test di Caserta.....	43
Fig. 3.5	Toy network: risultati sperimentali con perturbazione della stima a priori di tipo uniforme e 4 sezioni di conteggio	47
Fig. 3.6	Toy network: risultati sperimentali con perturbazione della stima a priori di tipo uniforme e 8 sezioni di conteggio	48
Fig. 3.7	Toy network: risultati sperimentali con perturbazione della stima a priori di tipo normale e 4 sezioni di conteggio.....	49
Fig. 3.8	Toy network: risultati sperimentali con perturbazione della stima a priori di tipo uniforme e 4 sezioni di conteggio	50
Fig. 3.9	Rete di Caserta: risultati sperimentali con perturbazione della stima a priori di tipo uniforme e 25 sezioni di conteggio	52
Fig. 3.10	Rete di Caserta: risultati sperimentali con perturbazione della stima a priori di tipo uniforme e 43 sezioni di conteggio	53
Fig. 3.11	Rete di Caserta: risultati sperimentali con perturbazione della stima a priori di tipo normale e 25 sezioni di conteggio.....	54
Fig. 3.12	Rete di Caserta: risultati sperimentali con perturbazione della stima a priori di tipo normale e 43 sezioni di conteggio.....	55
Fig. 4.1	Schema della metodologia proposta	59
Fig. 4.2	Rete di Torino,	60

Fig. 4.3	cvRMSE hold out flussi d'arco (setting 1) al variare delle stime a priori della domanda, senza introduzione di un errore in fase di assegnazione.....	66
Fig. 4.4	cvRMSE hold out flussi d'arco (setting 1) al variare delle stime a priori della domanda, con introduzione di un errore in fase di assegnazione.....	67
Fig. 4.5	cvRMSE hold out flussi d'arco (Setting 2) al variare del tasso di campionamento	69
Fig. 5.1	cvRMSE hold out flussi d'arco, $q=\rho=b=\gamma=0$, $\Phi=90k$, sezioni di conteggio=300, stima a priori da modello.....	81
Fig. 5.2	cvRMSE hold out flussi d'arco, $q=\rho=b=\gamma=0$, $\Phi=(5k\div 45k)$, sezioni di conteggio=300, stima a priori da modello.....	83
Fig. 5.3	cvRMSE hold out flussi d'arco, $q=\rho=b=\gamma=0$, $\Phi=90k$, sezioni di conteggio=50÷700, stima a priori da modello.....	85
Fig. 5.4	cvRMSE hold out flussi d'arco, $q=\rho=b=0$, $\gamma=50$, $\Phi=90k$, sezioni di conteggio: 300, stima a priori da modello.....	87
Fig. 5.5	cvRMSE hold out flussi d'arco, $q=\rho=b=0$, $\gamma=100$, $\Phi=90k$, sezioni di conteggio= 300, stima a priori da modello.....	89
Fig. 5.6	cvRMSE hold out flussi d'arco, $q=\rho=b=0$, $\gamma=50$, $\Phi=5k\div 45k$, sezioni di conteggio= 300, stima a priori da modello.....	90
Fig. 5.7	cvRMSE hold out flussi d'arco, $q=\rho=b=0$, $\gamma=100$, $\Phi=(5k\div 45k)$, sezioni di conteggio= 300, stima a priori da modello.....	92
Fig. 5.8	cvRMSE hold out flussi d'arco, $q=\rho=b=\gamma=0$, $\Phi=90k$, sezioni di conteggio=300, stima a priori di tipo diretta.....	93
Fig. 5.9	cvRMSE hold out flussi d'arco, $q=\rho=b=\gamma=0$, $\Phi=5k-45k$, sezioni di conteggio=300, stima a priori di tipo diretta $\alpha=1\%$	95
Fig. 5.10	cvRMSE hold out flussi d'arco, $q=\rho=b=\gamma=0$, $\Phi=5k-45k$, sezioni di conteggio=300, stima a priori di tipo diretta $\alpha=2\%$	96
Fig. 5.11	cvRMSE hold out flussi d'arco, $q=\rho=b=\gamma=0$, $\Phi=5k-45k$, sezioni di conteggio=300, stima a priori di tipo diretta $\alpha=3\%$	97
Fig. 5.12	cvRMSE hold out flussi d'arco, $q=\rho=b=\gamma=0$, $\Phi=5k-45k$, sezioni di conteggio=300, stima a priori di tipo diretta $\alpha=4\%$	98

Fig. 5.13	cvRMSE hold out flussi d'arco, $q=\rho=b=\gamma=0$, $\Phi=4k-45k$, sezioni di conteggio=300, stima a priori di tipo diretta $\alpha=5\%$	99
Fig. 5.14	cvRMSE hold out flussi d'arco, $q=\rho=b=\gamma=0$, $\Phi=5k-45k$, sezioni di conteggio=300, stima a priori di tipo diretta $\alpha=6\%$	100
Fig. 5.15	cvRMSE hold out flussi d'arco, $q=\rho=b=\gamma=0$, $\Phi=5k-45k$, sezioni di conteggio=300, stima a priori di tipo diretta $\alpha=7\%$	101
Fig. 5.16	cvRMSE hold out flussi d'arco, $q=\rho=b=\gamma=0$, $\Phi=5k-45k$, sezioni di conteggio=300, stima a priori di tipo diretta $\alpha=8\%$	102
Fig. 5.17	cvRMSE hold out flussi d'arco, $q=\rho=b=\gamma=0$, $\Phi=5k-45k$, sezioni di conteggio=300, stima a priori di tipo diretta $\alpha=9\%$	103
Fig. 5.18	cvRMSE hold out flussi d'arco, $q=\rho=b=\gamma=0$, $\Phi=5k-45k$, sezioni di conteggio=300, stima a priori di tipo diretta $\alpha=10\%$	104
Fig. 5.19	cvRMSE hold out flussi d'arco, $q=(0\div 1), \rho=b=\gamma=0$, coppie od variabili $\Phi=90k$, sezioni di conteggio=300, stima a priori da modello	106

1. Introduzione

1.1. Background

Ai fini della modellazione del funzionamento di un sistema di trasporto è essenziale stimare la domanda di mobilità, definita come la somma degli spostamenti di individui, veicoli o merci per un certo motivo in un dato intervallo temporale. La stima della domanda di mobilità è caratterizzata da una serie di complessità intrinseche, dovute al fatto che la domanda è generata da scelte di decisori (passeggeri, operatori nel campo delle merci, ecc.) e che, dunque, è funzione di complessi meccanismi comportamentali. Tradizionalmente, la stima della domanda è stata ottenuta con due differenti approcci: stima diretta (inferenza su dati campionari relativi a spostamenti e funzione di scelte individuali) e stime da modelli analitici (descrittivi, di regressione, comportamentali, etc.).

Sin dagli anni '70, è andato crescendo l'interesse per metodi di stima della domanda alternativi agli approcci tradizionali appena descritti. In particolare, data la semplicità e l'economicità di ottenimento di misure aggregate di traffico (es. conteggi manuali nelle sezioni stradali, dati provenienti da dispositivi di rivelazione automatica etc.), è andato affermandosi l'interesse verso metodologie di stima che prevedano tali misure di traffico come input.

Di converso, il problema di stima dei flussi origine-destinazione (da qui in poi, definiti semplicemente come flussi *od*) da misure di traffico, per ciascun modo di trasporto, rappresenta un problema matematico notevolmente complesso, in quanto fortemente indeterminato, a causa del fatto che il numero di flussi *od* incogniti è significativamente superiore al numero di misure generalmente disponibili, le quali rappresentano le equazioni del sistema. A titolo esemplificativo, per fornire un ordine di grandezza del problema in esame, si considera la situazione realistica relativa alla rete di Torino analizzata nel

successivo cap. 5, nella quale si è assunta la disponibilità di circa 300¹ misure di flussi di traffico, provenienti da altrettante sezioni di conteggio (equazioni), e in cui il numero di flussi *od* da stimare risulta superiore a 90'000 (incognite).

In letteratura scientifica, sono riconoscibili tre distinti macro-approcci per la risoluzione del problema dello sbilancio tra equazioni e incognite:

- 1) Utilizzo di una stima a priori della domanda: secondo tale approccio, la stima della domanda (ottenuta precedentemente mediante stima diretta o stima da modelli analitici) rappresenta la stima a priori, o non condizionata, mentre la conseguente stima condizionata alla stima a priori e alle informazioni di traffico aggregate rappresenta la stima a posteriori, o condizionata alle osservazioni disponibili. In particolare, dato che un contesto indeterminato è caratterizzato da infinite soluzioni al problema (vettori di domanda di trasporto) compatibili con i vincoli imposti dalle sezioni di misura, l'utilizzo di una stima a priori del vettore di domanda consente di ridurre la variabilità della nuova stima del vettore di domanda ricercato. Tale procedura prende il nome di *correzione della domanda di trasporto sulla base di misure aggregate di traffico*. Ovviamente, all'interno di tale procedura, la qualità della stima a priori della domanda risulta essere un input sensibile per la procedura, in quanto introduce potenzialmente delle distorsioni nelle stime a posteriori ricavate.
- 2) Incremento del numero di equazioni: rappresenta una strada alternativa o ad integrazione della precedente. Tale approccio prevede che si assuma l'ipotesi di dinamica intra-periodale (o *within day*) per la domanda di mobilità, discretizzando il periodo di analisi in intervalli temporali più piccoli, detti *time slice*. In tale contesto, ogni sezione di misura, per ciascuna delle time slice, rappresenta un'equazione del sistema. La metodologia appena descritta, però, comporta anche un incremento del numero di incognite, rappresentate dai vettori di domanda in corrispondenza di ciascuna time slice. In tal caso, assumendo delle ipotesi restrittive sull'evoluzione della domanda nel tempo, è possibile limitare l'incremento del numero di incognite rispetto all'incremento del numero di equazioni, ottenendo una complessiva riduzione dello sbilancio tra equazioni e incognite.

¹Il numero di sezioni di conteggio realmente installate nella città di Torino è circa 300.

3) Riduzione del numero di incognite: secondo tale approccio, si ipotizza che il vettore di domanda sia descrivibile attraverso un numero di variabili inferiore al numero di coppie *od*. In tal caso, in generale, si passa da una trattazione del problema in uno spazio con un numero molto elevato di dimensioni (il numero di coppie *od*) a uno spazio di dimensioni ridotte, introducendo un errore di approssimazione.

Il presente elaborato di tesi analizza alcuni avanzamenti metodologici relativi al terzo macro-approccio, discutendo, quando necessario, anche la importanza della disponibilità di una corretta stima a priori della domanda.

1.2. Outline

La presente tesi di dottorato ha l'obiettivo di migliorare le metodologie di stima della domanda sulla base di misure aggregate, con particolare riferimento alle tecniche per la riduzione dimensionale delle incognite.

Il successivo capitolo 2 presenta sinteticamente il problema della stima della domanda, ivi compreso quello della stima tramite conteggi di traffico, presentando la relativa letteratura scientifica rilevante. Nel corso del lavoro di tesi, sono state sviluppate tre diverse metodologie finalizzate alla riduzione dimensionale delle incognite. Per ciascuna delle metodologie proposte è stata considerata la presenza di una stima a priori della domanda. Pertanto, da questo punto in poi, verrà fatto riferimento alle tecniche proposte come metodologie alternative alla procedura di correzione della domanda di trasporto attraverso misure aggregate di traffico.

La prima metodologia proposta è un'euristica basata su un partizionamento (clustering) sequenziale di coppie *od*, mirato ad ottenere un sistema con un ridotto numero di incognite (le partizioni di coppie *od*). L'aggregazione delle coppie *od* in partizioni viene effettuata attraverso due criteri che verranno approfonditi nel cap. 3: adiacenza topologica delle zone e copertura dei flussi *od* (*coverage*). L'applicazione della metodologia proposta ha consentito di investigare le potenzialità della procedura di correzione della domanda al decrescere del numero di incognite (o, in altri termini, all'aumentare del livello di aggregazione), evidenziando dei significativi miglioramenti rispetto alle procedure di correzione classiche.

Successivamente, sulla base dei risultati ottenuti con la precedente procedura, è stata sviluppata una procedura di aggregazione che supera la logica di tipo sequenziale, e che si basa su un clustering dei flussi *od* effettuato in maniera simultanea. Il vantaggio di questa seconda metodologia, con riferimento alla prima, risiede nel significativo incremento di efficienza computazionale, che consente di estendere l'applicabilità della procedura di aggregazione dei flussi *od* anche a reti di grandi dimensioni. I dettagli di tale metodologia e le sperimentazioni condotte sono riportati nel cap. 4 del presente elaborato.

Infine, l'ultima procedura sviluppata si basa sulla riduzione dimensionale della stima dei flussi *od* attraverso un modello di tipo gravitazionale generalizzato, che adopera una tipologia di algebra deformata detta *q*-algebra e che consente di introdurre dei processi autoregressivi in grado di cogliere le correlazioni spaziali tra flussi *od*. Tale procedura consente la generalizzazione di metodologie già esistenti in letteratura, quali i modelli gravitazionali di tipo log-lineare semplice e una procedura analoga all'analisi fattoriale. La formalizzazione matematica della procedura e i risultati su una rete test reale sono riportati nel cap. 5 del presente elaborato di tesi.

Il capitolo 6 dell'elaborato riporta le principali conclusioni e suggerisce spunti per estensioni ulteriori sul lavoro di ricerca sin qui condotto.

2. Stima della domanda: generalità

Nel presente capitolo viene riportata un'analisi di letteratura relativa alle metodologie di stima della domanda di trasporto. Nel paragrafo 2.1 sono descritte le operazioni necessarie per lo studio di un sistema di trasporto preliminari alla stima della domanda. Nei paragrafi 2.2, 2.3 e 2.4, vengono descritti, rispettivamente, i seguenti tre metodi di stima della domanda di trasporto: stima diretta, stima da modelli e stima dei flussi di domanda da conteggi di traffico.

2.1. Studio di un sistema di trasporto

L'ingegneria dei trasporti è una branca dell'ingegneria civile che si occupa di studiare i sistemi di trasporto intesi come l'insieme di elementi fisici ed organizzativi che interagiscono per creare opportunità di trasporto e domanda di mobilità (Cascetta, 2006).

I modelli matematici utilizzati per lo studio dei sistemi di trasporto possono essere suddivisi in modelli di offerta e di domanda, i quali risultano caratterizzati da una interazione mutua. Pertanto, è necessario studiare anche l'interazione tra domanda ed offerta attraverso una procedura definita assegnazione.

Per studiare un sistema di trasporto è necessario individuare un'area di studio. Quest'ultima può essere definita come l'area all'interno della quale si ritiene che si esauriscano la maggior parte degli impatti relativi agli interventi di mobilità oggetto di analisi. Tale area viene separata dall'ambiente esterno attraverso una linea immaginaria, detta cordone. L'area di studio, per sua natura definita su un continuo spaziale, viene successivamente suddivisa in un insieme finito zone, attraverso una procedura definita zonizzazione. Le singole zone vengono individuate mediante il principio dell'omogeneità

trasportistica (Fig.1.1), definendo, per ciascuna di esse, un punto rappresentativo (generalmente in posizione baricentrica) detto *centroide interno*. I *centroidi esterni*, invece, vengono posizionati all'intersezione tra le infrastrutture rilevanti della rete di trasporto ed il cordone dell'area di studio. In generale, i centroidi rappresentano dei punti fittizi associati alle zone da cui si immagina, nella rappresentazione del modello, che partano o che arrivino tutti gli spostamenti da o verso ciascuna zona dell'area di studio.

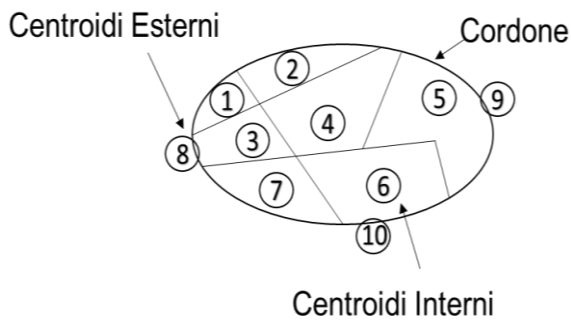


Fig. 1.1 Zonizzazione

o\d	1	2	3	4	5	6	7	8	9	10
1										
2										
3										
4				I					I-E	
5										
6										
7										
8								E-I		
9										
10									E-E	

Fig. 1.2 Matrice od

La domanda di trasporto rappresenta il numero di veicoli, persone o merci, che si spostano tra le zone di omogenee di traffico incluse in un'area di studio ed è un input fondamentale per qualsiasi analisi di tipo trasportistica. Una rappresentazione comunemente adoperata per i flussi di domanda è basata sulla cosiddetta matrice origine-destinazione (da qui in poi, definita semplicemente come matrice *od* o matrice di domanda; Fig.1.2). Le righe e le colonne di tale matrice rappresentano, rispettivamente, i centroidi di origine e di destinazione degli spostamenti, mentre il singolo elemento della matrice identifica il numero medio di spostamenti (*flussi od*) diretti dalla corrispondente origine alla corrispondente destinazione.

La matrice di domanda Fig.1.2, a sua volta, può essere suddivisa in quattro matrici:

- Matrice di mobilità Interna (I): essa rappresenta la mobilità interna all'area di studio, in quanto, sia i centroidi di origine, sia quelli di destinazione, sono interni all'area di studio.
- Matrice di scambio esterno-interno (E-I): in tale matrice le origini degli spostamenti costituiscono centroidi esterni all'area di studio, mentre le destinazioni costituiscono centroidi interni all'area di studio.

- Matrice di scambio interno-esterno (I-E): in tale matrice le origini degli spostamenti costituiscono centroidi interni all'area di studio, mentre le destinazioni costituiscono centroidi esterni all'area di studio.
- Matrice di attraversamento (E-E): in tale matrice sia i centroidi di origine che di destinazione sono esterni all'area di studio.

La stima della domanda di trasporto può essere condotta mediante due approcci distinti: stima diretta (mediante dati raccolti da indagini campionarie) e applicazione di modelli analitici.

Nel caso di matrici di scambio e di attraversamento la stima può avvenire esclusivamente per via diretta (indagini campionarie), mentre la matrice di mobilità interna può essere stimata anche attraverso l'applicazione di modelli analitici.

Un modello analitico per la stima della domanda di mobilità (da qui in poi richiamato sinteticamente come *modello di domanda*) è rappresentato dall'insieme delle relazioni matematiche che consentono di stimare il flusso di domanda di persone, veicoli o merci, che si spostano tra diversi punti di un sistema di trasporto, in un certo intervallo temporale, per un dato motivo, con una certa modalità di trasporto, attraversando una certa sequenza di archi della rete.

Nel caso di modelli di domanda riferiti ai passeggeri, è possibile specificare dei modelli analitici per caratterizzare in maniera differente la domanda di mobilità:

- Modelli per la stima dei flussi *od*;
- Modelli per la ripartizione modale;
- Modelli di scelta del percorso (tipicamente adoperati per le reti di trasporto privato o per le reti di trasporto collettivo con approcci a corse) o dell'ipercammino (tipicamente adoperati per le reti di trasporto collettivo con approcci a frequenze)².

Nel caso di trasporto merci, invece, vi sono ulteriori sotto-modelli in cui è possibile suddividere i modelli di domanda, quali:

- Modelli per la stima delle matrici di produzione-consumo;
- Modelli per la stima dei flussi *od*;

² La suddivisione riportata si riferisce alla struttura analitica comunemente definita modello a 4 stadi. Vi è la possibilità di caratterizzare ulteriormente la domanda di mobilità (es. tipo di sosta, sequenza di modi, orario di partenza, possesso dell'autoveicolo etc).

- Modelli per la ripartizione modale;
- Modelli di scelta del percorso o dell'ipercammino.

I modelli per la stima delle matrici produzione-consumo hanno l'obiettivo di stimare i flussi merci di categoria k , spostati tra la zona di produzione e la zona di consumo, all'interno dell'area di studio, in un dato intervallo temporale τ . Siccome tra il luogo di produzione e consumo sono interposti una serie di luoghi/fasi dello spostamento, le matrici produzioni-consumo inglobano a loro volta le matrici od , le quali sono associate agli spostamenti intermedi.

I modelli per la stima dei flussi od hanno l'obiettivo di stimare i flussi che si spostano tra ciascuna zona di origine o e zona di destinazione d all'interno dell'area di studio, in un dato intervallo temporale τ . Nel caso di trasporto passeggeri le zone di origine e destinazione coincidono con delle zone di traffico in cui è stata suddivisa l'area di studio. Nel caso di modelli per il trasporto merci, l'origine e la destinazione degli spostamenti sono rappresentati da luoghi intermedi appartenenti alle catene logistiche (centri di distribuzione, hub intermodali, etc.).

I modelli di ripartizione modale hanno, invece, l'obiettivo di stimare i flussi di persone, veicoli o merci stimati mediante i modelli per la stima dei flussi od che viaggiano a bordo di ciascuna modalità di trasporto m .

I modelli di scelta del percorso hanno l'obiettivo di stimare i flussi che viaggiano su una certa coppia od , con una certa modalità di trasporto m attraverso ciascuna generica sequenza di archi identificata come percorso p . I modelli di scelta dell'ipercammino, analogamente, hanno l'obiettivo di stimare i flussi che viaggiano su una certa coppia od , una certa modalità di trasporto m considerando un certo insieme di percorsi p , definito come ipercammino h . In tal caso, la strategia di viaggio assunta per gli utenti è di tipo preventivo-adattivo, in quanto si assume che vi sia un insieme di scelte *preventive* che possono essere effettuate prima di cominciare lo spostamento (es. scelta della fermata) e scelte *adattive* di cui l'utente ha conoscenza solo durante lo spostamento (es. linee che transitano ad una fermata scelta e che sono utili a raggiungere la destinazione).

Contestualmente alla costruzione del modello di domanda, per lo studio di un sistema di trasporto, è necessario implementare anche un *modello di offerta*. Tale modello è una rappresentazione schematica e parziale del sistema reale oggetto di studio. Ad esempio, il

sistema “città” può essere schematizzato a livello topologico attraverso un grafo composto da nodi e archi orientati, i quali, rispettivamente, risultano rappresentativi delle intersezioni e dei diversi tratti di arterie stradali. Non necessariamente, però, un arco risulta essere associato ad un’infrastruttura fisica. Esistono, infatti, archi mediante i quali possono essere rappresentate le rotte marittime ed aeree, oppure archi che simulano il tempo di attesa alla fermata di un mezzo pubblico. Per arrivare alla definizione di rete di trasporto, è necessario associare ad ogni elemento del grafo delle funzioni di costo come il tempo di percorrenza dell’arco, il pedaggio o il tempo di attesa ad un nodo. In altre parole, oltre al modello topologico, per rappresentare una rete di trasporto, risulta necessario associare anche un modello analitico, che consiste nell’insieme delle relazioni matematiche atte ad individuare le prestazioni identificative del livello del servizio di ciascun elemento della rete (impedenze, tempi di percorrenza, funzioni di utilità o di costo generalizzato di trasporto, etc.).

Una volta stimata la domanda di trasporto e definito il modello di offerta, è necessario trasformare gli elementi della matrice *od* in flussi sui singoli archi della rete mediante una procedura definita assegnazione. Tale procedura mette a sistema i modelli di domanda e di offerta. Esistono differenti procedure di assegnazione, principalmente in funzione delle ipotesi alla base dei modelli di scelta del percorso.

I flussi di arco da modello risultano affetti da una distorsione rispetto ai flussi veri, in virtù del fatto che i modelli analitici introducono una serie di ipotesi semplificative rispetto ai complessi fenomeni comportamentali soggiacenti le scelte che determinano la domanda di mobilità. Ad ogni modo, data la semplicità di reperimento di misure aggregate del traffico (es. conteggi alle intersezioni stradali), queste ultime possono essere adoperate per ottenere delle stime della domanda corrette rispetto a quelle originarie. Questa procedura prende il nome di *“correzione della domanda di trasporto sulla base di dati aggregati di traffico”*. Tale procedura di correzione si basa sulla risoluzione di un sistema di equazioni stocastiche in cui le equazioni sono rappresentate dai flussi d’arco osservati in corrispondenza delle sezioni di conteggio della rete, mentre le incognite sono gli elementi della matrice *od*. Generalmente nelle reti reali il numero di sezioni di conteggio disponibili risulta di gran lunga inferiore rispetto al numero di coppie *od*, rendendo così il problema fortemente indeterminato. Inoltre, tra le infinite possibili soluzioni compatibili con i vincoli imposti

dalle sezioni di conteggio, in genere si sceglie quella che si avvicina di più alla stima a priori della domanda. Studi di letteratura hanno evidenziato come l'efficacia della procedura di correzione della matrice *od* migliori man mano che il rapporto tra equazioni ed incognite diviene prossimo all'unità. Per ottenere la condizione di bilancio equazioni-incognite, finalizzata al raggiungimento di una stima che sia più affidabile di quella restituita dalle classiche procedure di correzione, è possibile perseguire due approcci. Un primo approccio consiste nell'incrementare il numero di equazioni, mentre il secondo approccio si persegue tentando di ridurre il numero di incognite. Entrambe le procedure, però, non consentono di risolvere in modo definitivo il problema della correzione *od*. Vantaggi e svantaggi di tali procedure saranno descritti nei prossimi paragrafi.

2.2. Stima della domanda di mobilità: stima diretta

L'obiettivo di un'indagine campionaria per la stima diretta della domanda è ottenere una stima della variabile "domanda" a livello dell'intera popolazione a partire da osservazioni su un sottoinsieme di dati definito campione. Trattazioni sull'argomento sono riportate in diversi volumi (Cascetta, 2006; Ortúzar & Willumsen, 2011).

La domanda di trasporto nel caso delle matrici di scambio e attraversamento deve essere necessariamente stimata attraverso delle indagini per la stima diretta. Le indagini possono essere condotte durante il viaggio "*a bordo*", cioè si intervista un campione di utenti di uno o più modi di trasporto. Le interviste condotte a bordo strada possono riguardare il conducente e i passeggeri nel caso di trasporto privato (Simonelli et al., 2020), gli autisti nel caso di trasporto merci, oppure, possono essere condotte all'interno dei mezzi di trasporto o nelle stazioni/terminal nel caso di trasporto pubblico.

Per quanto riguarda la stima della mobilità interna, invece, è possibile utilizzare delle indagini definite a *domicilio* in cui si intervista un campione di persone, famiglie, operatori logistici all'interno dell'area di studio. Le indagini a domicilio hanno un'elevata affidabilità dovuta all'interazione fisica tra l'intervistato e l'intervistante, ma per essere efficaci hanno bisogno di un numero ingente di interviste, di conseguenza hanno costi elevati. In alternativa all'intervista a domicilio eseguita di persona, esistono altre tipologie di interviste condotte a telefono o mediante piattaforme web, ma a fronte di un notevole risparmio economico forniscono risultati meno affidabili.

La progettazione di un'indagine campionaria richiede la definizione di:

- unità di campionamento;
- strategia di campionamento;
- stimatore da adottare;
- numerosità del campione.

L'unità di campionamento rappresenta l'unità base dell'intervista. Nel caso di interviste a bordo strada, l'unità può essere rappresentata dal veicolo oppure dai singoli utenti in attesa alla fermata del mezzo pubblico.

La strategia di campionamento utilizzata è spesso di tipo probabilistico e consente di definire a priori dei possibili risultati in funzione della strategia adottata. Il campione è estratto in modo da garantire un prefissato risultato. Tra le strategie di campionamento più diffuse vi sono:

- campionamento casuale semplice: tutti gli elementi della popolazione hanno uguale probabilità di appartenere al campione estratto;
- campionamento casuale stratificato: la popolazione è divisa in gruppi detti *strati*, non sovrapposti ed esaustivi, da ciascuno strato è estratto un campione casuale semplice. Gli elementi di ciascuno strato hanno stessa probabilità di appartenere al campione, mentre a elementi di strati differenti, corrispondono probabilità di appartenere al campione diverse.
- Campionamento a grappolo: le unità di riferimento sono raccolte in grappoli (famiglie di un condominio, passeggeri di un veicolo) i quali, nel caso di campionamento semplice, vengono estratti con una prefissata probabilità di appartenere al campione, oppure, vengono estratti con probabilità differenti nei vari strati nel caso di campionamento stratificato.

Lo stimatore da adottare è funzione della strategia di campionamento utilizzata.

La numerosità campionaria dipende dalla varianza che si vuole fissare per la stima; minore è la varianza della stima, maggiore sarà la numerosità del campione.

Un'altra modalità con cui è possibile eseguire la stima diretta della domanda è attraverso l'utilizzo delle traiettorie GPS. Tali tecniche hanno avuto un notevole sviluppo nell'ultimo decennio grazie alla presenza sempre maggiore di dispositivi tecnologici a bordo dei veicoli

quali, scatole nere, navigatori, antifurti satellitari, applicazioni per smartphone (Google Maps, Waze, etc.). Anche gli operatori telefonici, grazie alla triangolazione delle celle GSM, sono in grado fornire informazioni relative alle traiettorie. Infine, la tecnologia bluetooth presente in smartphone ed autoveicoli può essere utilizzata per la stima della domanda di trasporto (Barceló et al., 2013). Data la grande di disponibilità dei dati appena citati, si è sviluppato un nuovo segmento di mercato relativo al commercio di dati di mobilità, con conseguente utilizzo dei precedenti per analisi trasportistiche. *INRIX*, *OCTOTELEMATICS* e *VODAFONE* sono solo alcune delle innumerevoli società dedite alla raccolta e alla vendita di traiettorie GPS.

I dati raccolti possono essere sfruttati in diversi modi, tra cui stima diretta della matrice di domanda; stima della matrice di assegnazione; stima dei tassi di diversione alle intersezioni. Esempi di utilizzo di traiettorie per la stima e la correzione della domanda sono riportati in diversi articoli scientifici (Gundlegard & Karlsson, 2009; Wang et al., 2010; Barceló et al., 2013; Demissie et al., 2013; Iqbal et al., 2014; Park et al., 2014; Larijani et al., 2015; Ge & Fukuda, 2016; Huang et al., 2019)

Ovviamente l'affidabilità dei risultati è legata al tasso di campionamento delle traiettorie (Simonelli et al., 2019), definito come il rapporto tra il numero di spostamenti campionati rispetto al totale. Sfortunatamente, nei casi reali, i tassi di campionamento sono spesso troppo bassi per poter eseguire delle buone stime, oltre al fatto che i dati hanno dei costi di acquisto notevoli. Dunque, è sempre preferibile applicare una procedura di correzione mediante utilizzo di misure aggregate anche ad una stima diretta della domanda ottenuta tramite l'uso di dati di traiettorie.

Nel caso del trasporto merci, invece, la stima diretta della domanda può avvenire attraverso la consultazione di diversi database disponibili on-line che spesso forniscono dati gratuiti a diversa scala geografica per le diverse categorie merceologiche e modalità di trasporto utilizzate. Esempi di questi database sono disponibili nei portali web di *UNCTAD*, *COMTRADE*, *CEPII*, *EUROSTAT*.

2.3. Stima della domanda di mobilità: modelli analitici

La matrice di domanda può essere stimata applicando un sistema di modelli, il cui risultato è il numero di spostamenti medio tra un'origine o e destinazione d in una certa fascia temporale τ , per un certo motivo di spostamento s , per utenti di una certa categoria c . Tali modelli, per poter essere applicati, necessitano della stima di parametri sulla base di osservazioni reali (calibrazione). Risulta necessario, dunque, anche in questo caso, procedere a delle indagini per la raccolta di dati reali preliminarmente alla fase di applicazione dei modelli, ovviamente, il numero di parametri da stimare è molto ridotto rispetto ad un'indagine per la stima diretta della domanda.

In generale un modello può essere definito come una funzione matematica del tipo $Y = f(X)$ finalizzata a riprodurre una relazione causa-effetto, dove Y rappresenta la domanda da stimare funzione di una serie di variabili esplicative X .

Un primo insieme di variabili da cui si immagina dipendere la domanda di trasporto nei modelli di domanda, sono le variabili di tipo socioeconomiche **SE**.

$$d_{od} = d_{od}(SE) \quad (2.1)$$

Il principio su cui si basano i modello di domanda è che una zona genera flussi proporzionalmente alla propria massa, che in funzione dei casi (passeggeri/merci), può essere espressa in termini di residenti, addetti, numero di industrie, PIL o PIL pro capite, etc. Anche la capacità di attrarre un flusso da parte di una zona può essere messa in relazione ad una massa di quella zona, anche in questo caso variabile a seconda dei casi (numero di addetti, di residenti, di addetti alle industrie di trasformazione dei beni, densità di servizi ed attività commerciali, industrie, etc). I precedenti dati sono in genere disponibili presso i portali nazionali che si occupano di statistica, ad esempio, in Italia le informazioni sopracitate sono reperibili gratuitamente dal sito dell'*ISTAT*.

L'altra tipologia di variabile da cui si immagina dipendere la domanda nei modelli è l'impedenza di trasporto **T**.

$$d_{od} = d_{od}(SE, T) \quad (2.2)$$

Secondo tali modelli una coppia di zone tenderà a scambiare flussi con una relazione di proporzionalità inversa all'impedenza di trasporto che le separa. L'impedenza di trasporto può essere valutata in diversi modi, ad esempio, come distanza, tempo di percorrenza sulla

rete, costo generalizzato (somma omogeneizzata di tempi e costi) o semplicemente come costo di trasporto.

I modelli utilizzati per la stima dei flussi od^3 possono essere suddivisi in due macrocategorie:

- basati sull'utilità aleatoria
- descrittivi

I modelli basati sull'utilità aleatoria stimano il flusso di domanda d attraverso l'applicazione di due modelli in sequenza. Nel caso di modelli basati sull'utilità aleatoria di tipo "supply driven", i modelli in sequenza sono di emissione e di distribuzione, mentre nel caso di modelli basati sull'utilità aleatoria di tipo "demand driven", i modelli in sequenza sono di attrazione e di acquisizione.

I modelli descrittivi, invece, sono modelli in cui la relazione causa effetto risulta essere molto semplificata, un esempio, è il modello gravitazionale ispirato proprio alla legge di gravitazione universale di Newton.

A valle dei precedenti modelli, è possibile applicare ulteriori modelli (ripartizione modale, scelta del percorso)

I modelli di domanda di seguito descritti saranno specificati nel caso di trasporto delle merci, ma la trattazione può essere estesa ai modelli per il trasporto di passeggeri.

³ Storicamente si parla di flussi od anche nel caso delle merci, in realtà, l'applicazione di un modello di domanda nel caso di trasporto merci fornisce come risultato delle matrici produzione-consumo.

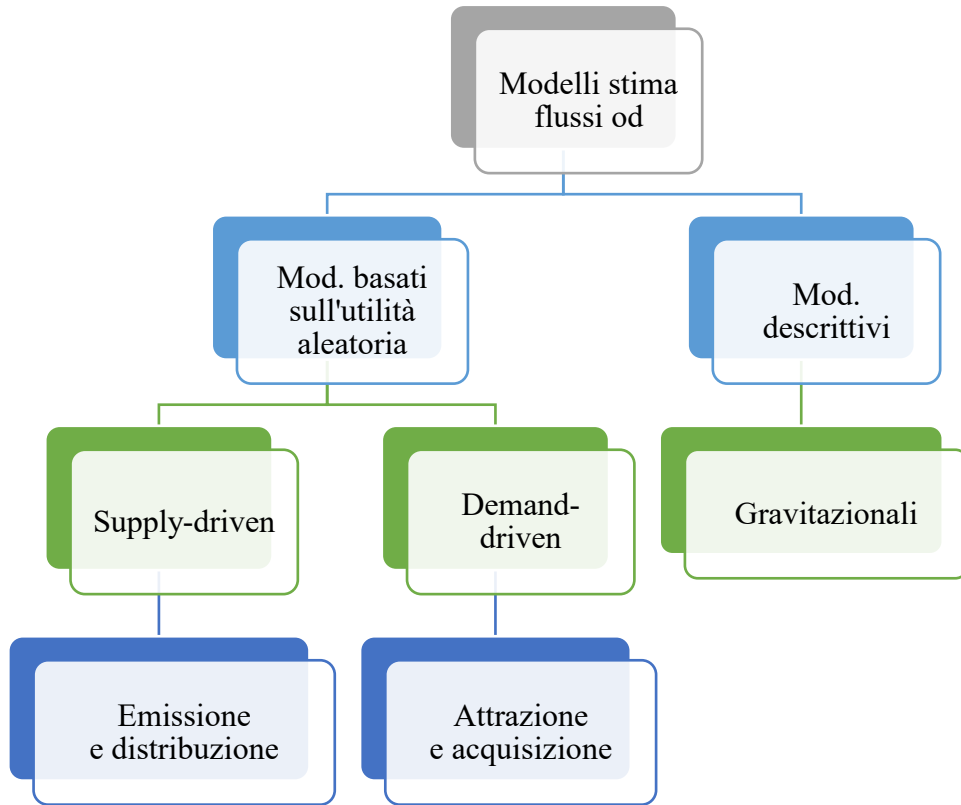


Fig. 2.1 Classificazione dei modelli di domanda

2.3.1. Modelli basati sull'utilità aleatoria

I modelli basati sull'utilità aleatoria sono modelli che consentono di stimare la domanda di trasporto attraverso l'applicazione di modelli in sequenza. A loro volta questi modelli possono essere suddivisi a seconda della grandezza che si immagina governare il mercato, infatti, si parla di modelli guidati dalla produzione o dai consumi. Della prima tipologia di modelli fanno parte i modelli supply-driven (guidati dalla produzione), in cui la grandezza da cui si immagina dipendere la domanda di trasporto è la produzione di beni.

$$d_{od,k} = d_{o,k} \cdot p[d|o, k] \quad (2.3)$$

$d_{o,k}$ rappresenta la domanda emessa dalla zona o relativamente alla categoria merceologica k e stimabile attraverso il modello di emissione (2.4).

$$d_{o,k} = \beta_e^k \cdot M_{o,k} \quad (2.4)$$

β_e^k è l'indice di emissione, costante per ogni zona, variabile per categoria merceologica k , e stimato attraverso una procedura di calibrazione. Il modello di emissione così come specificato nell'equazione 2.4 è definito *indice per categoria*.

$M_{o,k}$ rappresenta la massa nella zona di origine/produzione delle merci di tipologia k e può essere espressa come PIL, PIL pro capite o numero di addetti alle unità locali relativamente alla categoria merceologica k .

Una volta definito il flusso complessivo uscente dalla zona o di categoria k , è necessario ripartirlo verso le diverse zone di destinazione, cioè eseguire una disaggregazione spaziale mediante l'applicazione del modello di distribuzione (2.5).

$$p[d|o, k] = \frac{\exp\left(\frac{V_{d|o,k}}{\theta}\right)}{\sum_{d'} \exp\left(\frac{V_{d'|o,k}}{\theta}\right)} \quad (2.5)$$

$p[d|o, k]$ rappresenta la percentuale del flusso emesso dalla zona o relativamente alla categoria merceologica k , diretto nella zona di destinazione d . Il modello di distribuzione riportato nella (2.5) è, in particolare, il modello logit-multinomiale, in cui:

- θ rappresenta il parametro di varianza del modello;
- $V_{d|o,k}$ è l'utilità di esportare da o verso d i flussi merci di tipologia k ;

L'utilità di esportare in una zona è direttamente proporzionale ai consumi di quella zona e inversamente proporzionale all'impedenza di trasporto relativa alla relazione od ; una tipica specificazione adottata è la seguente:

$$V_{d|o,k} = \beta_t \cdot t_{od} + \beta_{c,k} \cdot C_{od,k} + \beta_{a,r}^k \cdot Residenti_d + \beta_{a,Imp}^k \cdot Addetti_{ULd}^k \quad (2.6)$$

L'espressione (2.6) è una somma omogeneizzata di diverse quantità attraverso dei coefficienti di omogeneizzazione β , la cui unità di misura è inversa a quella della grandezza moltiplicata. I consumi in destinazione d , dei beni di tipo k , dipendono dai residenti ($Residenti_d$) e dagli addetti alle unità locali delle industrie di trasformazione ($Addetti_{ULd}^k$).

t_{od} e $C_{od,k}$ sono rispettivamente tempi e costi necessari per spostare i beni di tipologia k tra o e d . Le due quantità rappresentano un'impedenza di trasporto.

β_i sono dei parametri del modello da calibrare.

La seconda tipologia di modelli basati sull'utilità aleatoria sono i modelli di tipo demand-driven (guidati dalla domanda), in cui la grandezza da cui si immagina dipendere la domanda di trasporto è il consumo di beni.

$$d_{od,k} = d_{d,k} \cdot p[o|d, k] \quad (2.7)$$

Secondo il modello $d_{d,k}$ rappresenta il flusso merci di tipologia k attratto dalla zona d , tale quantità può essere stimata mediante il modello di attrazione (2.7).

$$d_{d,k} = \beta_{a,r}^k \cdot Residenti_d + \beta_{a,Imp}^k \cdot Addetti_{ULd}^k \quad (2.8)$$

Il consumo di merci di tipo k dipende da due aliquote, residenti ed addetti alle unità locali cioè, lavoratori impiegati nelle industrie di trasformazione, in percentuali diverse e pari ai β .

$\beta_{a,r}^k$ e $\beta_{a,Imp}^k$ sono gli indici di attrazione relativamente a residenti ($Residenti_d$) e agli addetti ($Addetti_{ULd}^k$), costanti da zona a zona, variabili solo per categoria merceologica k .

Entrambi gli indici sono ottenuti attraverso un preventivo processo di calibrazione.

Una volta definito il flusso complessivamente attratto dalla zona d di categoria k , è necessario capire la provenienza, cioè eseguire una disaggregazione spaziale mediante l'applicazione del modello di acquisizione (2.9).

$$p[o|d,k] = \frac{\exp\left(\frac{V_{o|d,k}}{\theta}\right)}{\sum_{o'} \exp\left(\frac{V_{o'|d,k}}{\theta}\right)} \quad (2.9)$$

$p[o|d,k]$ rappresenta la percentuale del flusso importato nella zona d di tipologia k dalla zona o . Il modello di acquisizione 2.9 è ancora un modello logit-multinomiale, in cui:

- θ rappresenta il parametro di varianza;
- $V_{o|d,k}$ è l'utilità di importare nella zona d , merci di tipologia k dalla zona o ;

L'utilità è direttamente proporzionale alla produzione della zona di origine o e inversamente proporzionale all'impedenza di trasporto tra o e d ; infatti:

$$V_{o|d,k} = \beta_t \cdot t_{do} + \beta_{c,k} \cdot C_{do,k} + \beta_{e,o}^k \cdot Addetti_{ULO}^k \quad (2.10)$$

L'espressione 2.10 è una somma omogeneizzata di diverse quantità attraverso i coefficienti di omogeneizzazione β , la cui unità di misura è inversa a quella della grandezza moltiplicata. La produzione in origine o di merci di tipologia k dipende, ad esempio, dagli addetti alle unità locali ($Addetti_{ULO}^k$).

t_{od} e $C_{od,k}$ sono rispettivamente tempi e costi di trasporto tra o e d del bene k .

β_i sono dei parametri da calibrare.

2.3.2. Modelli descrittivi: modello gravitazionale

La principale caratteristica del modello gravitazionale, che ricade nella famiglia dei modelli descrittivi, è la capacità di stimare direttamente un flusso di domanda $d_{od,k}$, senza utilizzare modelli in sequenza, contrariamente a quanto avviene per i modelli basati sull'utilità aleatoria. Nel modello gravitazionale, si assume che due zone di traffico in un'area di studio possano scambiare un flusso di domanda proporzionale alle rispettive masse e inversamente proporzionale all'impedenza di trasporto che le separa. Un modello gravitazionale può essere specificato nel seguente modo:

$$d_{od,k} = \beta_{0,k} \cdot M_{o,k}^{\beta_{1,k}} \cdot M_{d,k}^{\beta_{2,k}} \cdot C_{od,k}^{*\beta_{3,k}} \cdot Ta_{od,k}^{\beta_{4,k}} \cdot \exp(\delta_{od}^{\beta_{5,k}}) \quad (2.11)$$

dove:

$M_{o,k}$ è la massa in origine rappresentata ad esempio dalla produzione e valutata come PIL, PIL pro capite o numero di addetti alle unità locali relativamente al settore di produzione della categoria merceologica k ;

$M_{d,k}$ è la massa in destinazione rappresentata dai consumatori. Può essere quantificata come residenti o addetti alle industrie di trasformazione;

$C_{od,k}^*$ è l'impedenza di trasporto tra o e d relativamente alla categoria merceologica k (costo generalizzato, distanza o tempo di percorrenza su rete, distanza in linea d'aria etc.);

$Ta_{od,k}$ è un dazio doganale o tariffa applicata per spostare le merci di tipologia k dalla zona o alla zona d ;

δ_{od} sono variabili di tipo dicotomiche [0/1], dette anche dummy. Possono essere riferite ad una singola zona (accesso al mare, condizione di conflitto bellico, etc.) o ad una coppia di zone (lingua comune, confine comune, accordo commerciale comune);

β sono dei parametri del modello da calibrare.

Dal punto di vista operativo si predilige una trasformazione logaritmica del modello gravitazionale in quanto più semplice da trattare in fase di calibrazione. Infatti, si utilizza una modello di regressione del tipo log-log e il modello diventa:

$$\begin{aligned} \ln(d_{od,k}) = & \ln(\beta_{0,k}) + \beta_{1,k} \cdot \ln(M_{o,k}) + \beta_{2,k} \cdot \ln(M_{d,k}) + \beta_{3,k} \cdot \ln(C_{od,k}) \\ & + \beta_{4,k} \cdot \ln(Ta_{od,k}) + \beta_{5,k} \cdot \delta_{od} \end{aligned} \quad (2.12)$$

In questa particolare specificazione il generico coefficiente β_X , associato all'attributo X , rappresenta la cosiddetta elasticità. Cioè, quanto varia la variabile dipendente a seguito di una variazione unitaria dell'attributo X

2.4. Correzione dei flussi *od* mediante conteggi di traffico

Un'altra strada che consente di stimare la domanda di trasporto è la correzione della domanda con conteggi di traffico. La procedura può essere eseguita sia in condizioni statiche sia dinamiche. Nel primo caso, si ipotizza che la domanda e l'offerta di trasporto abbiano caratteristiche costanti durante un periodo di tempo abbastanza lungo tale da raggiungere condizioni stazionarie, cioè, le variabili significative assumono valori indipendenti dal tempo. In condizioni dinamiche, invece, si simula come la domanda e l'offerta di trasporto varino all'interno del periodo di tempo considerato.

La procedura di correzione è anche definita come problema inverso dell'assegnazione, in quanto, l'assegnazione consente di trasformare i flussi del vettore di domanda $\mathbf{d}$ in flussi sui singoli archi della rete $\mathbf{f}$ (2.14).

$$\mathbf{f} = \mathbf{M} \cdot \mathbf{d} + \boldsymbol{\varepsilon}_{sim} \quad (2.13)$$

$\mathbf{M}$ rappresenta la matrice di assegnazione. Essa consente di assegnare, cioè, associare la domanda di trasporto $\mathbf{d}$ ai singoli archi della rete mediante una trasformazione lineare. In essa sono contenute tutte una serie di ipotesi sui modelli di scelta del percorso o sulle tecniche utilizzate per il campionamento da traiettorie. Di conseguenza $\mathbf{M}$ ingloba un errore di simulazione che può essere espresso attraverso il termine $\boldsymbol{\varepsilon}_{sim}$. Dunque, anche se la stima della domanda fosse perfetta, ci sarebbe un errore sui flussi d'arco dovuto alla matrice d'assegnazione.

Specularmente alla procedura di assegnazione, la procedura di correzione della domanda sfrutta le informazioni contenute nelle sezioni di conteggio per ricavare una stima della domanda di trasporto, e quindi identifica un nuovo vettore di domanda compatibile con i vincoli imposti dal problema.

I flussi di arco osservati nelle sezioni di conteggio, essendo ottenuti attraverso misurazioni (manuali, automatiche, ecc.), inglobano degli errori di misura $\boldsymbol{\varepsilon}_{obs}$. Dunque, la stima del flusso osservato $\hat{\mathbf{f}}$ sarà pari al flusso vero $\mathbf{f}$ più un errore di misura $\boldsymbol{\varepsilon}_{obs}$:

$$\hat{\mathbf{f}} = \mathbf{f} + \boldsymbol{\varepsilon}_{obs} \quad (2.14)$$

Combinando le equazioni (2.14) e (2.15), si ottiene una relazione che lega i flussi d'arco alla matrice di assegnazione e alla domanda di trasporto.

$$\begin{aligned} \hat{\mathbf{f}} = \mathbf{f} + \boldsymbol{\varepsilon}_{obs} &= \mathbf{M} \cdot \mathbf{d} + \boldsymbol{\varepsilon}_{sim} + \boldsymbol{\varepsilon}_{obs} = \mathbf{M} \cdot \mathbf{d} + \boldsymbol{\varepsilon} \\ \boldsymbol{\varepsilon} &= \boldsymbol{\varepsilon}_{obs} + \boldsymbol{\varepsilon}_{sim} \end{aligned} \quad (2.15)$$

La procedura di correzione ricerca quella particolare configurazione di $\mathbf{d}$ in grado di riprodurre $\hat{\mathbf{f}}$.

Le informazioni sulla domanda contenute nelle sezioni di conteggio, rappresentate dal sistema di equazioni stocastiche (2.16), non sono sufficienti a stimare il vettore $\mathbf{d}$. Infatti, anche assumendo $\boldsymbol{\varepsilon} = \mathbf{0}$, il numero di equazioni linearmente indipendenti sono in genere, nelle reti reali, in numero molto inferiore agli elementi del vettore di domanda $\mathbf{d}$ da stimare; quindi, si avrebbero infinite soluzioni al problema. Inoltre, essendo $\boldsymbol{\varepsilon}$ diverso dal vettore nullo, il sistema potrebbe non ammettere alcuna soluzione. Pertanto, per limitare la variabilità delle possibili soluzioni, è necessario aggiungere equazioni al sistema (2.16) provenienti da altre fonti (indagini campionarie oppure conoscenze a priori sui dati), in genere si utilizza una stima a priori della domanda di trasporto.

Contestualmente al problema dello sbilancio equazioni-incognite, vi è un altro fattore che influenza la qualità della procedura di correzione definito *od coverage* χ_{od} (2.17) che stabilisce il grado con il quale si riesce ad intercettare un flusso *od* sul complesso delle sezioni di conteggio possibili.

$$\chi_{od} = \sum_{c \in C} m_{od}^c \quad (2.16)$$

m_{od}^c rappresenta il generico elemento della matrice di assegnazione relativo alla sezione di conteggio c appartenente all'insieme delle sezioni di conteggio C , cioè la percentuale del flusso relativo alla relazione *od* che utilizza l'arco c .

Sulla definizione di *od coverage* χ_{od} si fonda un concetto molto importante, cioè, se un flusso *od* non è intercettato da alcuna sezione di conteggio non può essere corretto da qualsiasi procedura di correzione applicata. Di conseguenza è fondamentale identificare dei criteri di posizionamento ottimali delle sezioni di conteggio (Yang et al., 1991; Yang & Zhou, 1998; Bianco et al., 2001; Gan Liping et al., 2005; Ehlert et al., 2006; Yang et al., 2006; Chen et al., 2007; Simonelli et al., 2012).

Chiaramente le procedure citate sopra sono applicabili in contesti in cui non esistono già sezioni di conteggio o comunque in quelle situazioni in cui il gestore delle infrastrutture viarie è disposto ad investire nell'installazione di dispositivi di misurazione dei flussi di traffico. Però, nella stragrande maggioranza dei casi, la posizione delle sezioni di conteggio è nota e non modificabile; quindi, è la procedura di correzione a dover sfruttare al meglio tutte le informazioni disponibili provenienti dai conteggi.

La formulazione generale del problema della correzione della domanda passa attraverso un'ottimizzazione, che può essere espressa mediante la seguente relazione:

$$\mathbf{d}^* = \operatorname{argmin}_{\mathbf{x} > 0} [z_1(\mathbf{x}, \hat{\mathbf{d}}) + z_2(\mathbf{f}(\mathbf{x}), \hat{\mathbf{f}})] \quad (2.17)$$

dove:

$\hat{\mathbf{d}}$, $\mathbf{x}$ rappresentano rispettivamente una stima a priori della domanda di trasporto e la domanda incognita;

$\hat{\mathbf{f}}$, $\mathbf{f}(\mathbf{x})$ rappresentano rispettivamente la stima a priori dei flussi di arco, ottenuti mediante osservazione nelle sezioni di conteggio stradali, e i flussi di arco ottenuti assegnando la domanda incognita $\mathbf{x}$;

$z_1(\mathbf{x}, \hat{\mathbf{d}})$ è una funzione che misura la distanza tra la domanda incognita $\mathbf{x}$ e la sua stima a priori $\hat{\mathbf{d}}$;

$z_2(\mathbf{f}(\mathbf{x}), \hat{\mathbf{f}})$ è una funzione che misura la distanza tra i flussi d'arco ottenuti assegnando la domanda incognita $\mathbf{x}$ e gli omologhi osservati nelle sezioni di conteggio.

A seconda di come si specificano le funzioni di distanza z_1 e z_2 , è possibile ottenere diverse famiglie di stimatori (minimi quadrati generalizzati (GLS), massima verosimiglianza, stimatori bayesiani, etc.).

Un altro modo di eseguire una correzione della domanda di trasporto è quella di utilizzare le informazioni contenute nei conteggi di traffico per stimare/correggere i parametri dei modelli di domanda utilizzati per stimare la domanda stessa. In questo caso la domanda incognita è ottenuta attraverso l'applicazione di un modello di domanda che a sua volta è funzione di una serie di parametri da correggere. Ad esempio, i parametri potrebbero essere i $\boldsymbol{\beta}$ di un modello gravitazionale. In questo caso la procedura di ottimizzazione non lineare nei parametri potrebbe assumere la seguente espressione:

$$\boldsymbol{\beta}^* = \operatorname{argmin}_{\boldsymbol{\beta}} [z_3(\boldsymbol{\beta}, \hat{\boldsymbol{\beta}}) + z_4(\mathbf{f}(\boldsymbol{\beta}), \hat{\mathbf{f}})] \quad (2.18)$$

$$f(\beta) = M \cdot d(\beta) \quad (2.19)$$

dove:

$\beta, \hat{\beta}$ rappresentano i vettori di parametri da cui dipende il modello di domanda rispettivamente, vettore incognito e proveniente da una stima a priori (calibrazione disaggregata);

$z_3(\beta, \hat{\beta})$ rappresenta una misura di distanza tra i vettori di parametri β e $\hat{\beta}$

$z_4(f(\beta), \hat{f})$ rappresenta una misura di distanza tra i flussi d'arco osservati $\hat{f}$ e i flussi d'arco ottenuti assegnando la domanda generata da modello dipendente dai parametri incogniti β (2.20).

2.4.1. Approcci per ridurre lo sbilancio equazioni incognite

Lo sbilancio tra equazioni e incognite è affrontato in letteratura scientifica utilizzando due approcci diametralmente opposti:

Riduzione del numero di incognite, tale approccio consente di utilizzare un ridotto numero di incognite per spiegare la domanda di mobilità. Esempi di applicazioni in tal senso sono stati seguiti attraverso l'analisi per componenti principali *PCA*; (Djukic, Lint, et al., 2012; Djukic, Flötteröd, et al., 2012; Prakash et al., 2017, 2018). La metodologia basata sulla *PCA* consente di selezionare alcune componenti che spiegano una certa aliquota della varianza complessiva del problema trascurandone un'altra. Sebbene tale approccio sia metodologicamente valido, esso non ha trovato applicazione in contesti reali in quanto risulta di difficile applicazione a causa dell'elevato onere computazionale.

Altri approcci finalizzati alla riduzione dimensionale sono basati su tecniche di machine-learning.

Incremento del numero di equazioni, tale approccio assume una dinamica intra-periodale (o *within day*) per la domanda di mobilità, discretizzando il periodo di analisi in intervalli temporali più piccoli detti *time slice*. In questo modo ogni *time slice*, per ogni sezione di misura, sarà in grado di fornire un'equazione al sistema. Questo approccio, però, comporta anche un incremento del numero di incognite rappresentate da tanti vettori di domanda quanti gli intervalli temporali considerati. Però, attraverso l'introduzione di ipotesi sull'evoluzione della domanda tra le diverse *time slice*, è possibile limitare l'aumento dei

flussi *od* incogniti. Le ipotesi maggiormente utilizzate per spiegare l'evoluzione della domanda sono le seguenti:

- processi auto-regressivi (Balakrishna et al., 2005), mediante l'introduzione di una dipendenza lineare tra la domanda ad un certo istante e le domande negli istanti precedenti;

- l'ipotesi quasi-dinamica (Marzano et al., 2009; Cascetta et al., 2013; Cantelmo et al., 2015; Marzano et al., 2018), che consiste nell'assunzione di flussi generati da ciascuna origine variabili in ogni sezione temporale, e di probabilità di scelta della destinazione costanti per periodi più lunghi di tempo.

Tali approcci basati sull'analisi dinamica, si prestano meglio per applicazioni real-time o di controllo del traffico, piuttosto che per la pianificazione dei sistemi di trasporto.

Altre tecniche presenti in letteratura relativamente alla stima *od* sono basate su tecniche di soft-computing o swarm intelligence (Caggiani et al., 2012, 2013; Dantsuji et al., 2022; López-Ospina et al., 2021)

2.4.2. Approcci per garantire la coverage

Per garantire la condizione di coverage sarebbe necessario posizionare le sezioni di conteggio in modo tale da intercettare almeno una volta ogni flusso *od*. In letteratura vi sono diversi approcci descritti per il posizionamento ottimale le sezioni di conteggio (Yang et al., 1991; Yang & Zhou, 1998; Bianco et al., 2001; Gan Liping et al., 2005; Ehlert et al., 2006; Yang et al., 2006; Chen et al., 2007; Simonelli et al., 2012; Capar et al., 2013; Fei et al., 2013; He, 2013; Ng, 2013; Bianco et al., 2014), sicuramente uno dei più famosi e semplice da applicare è il criterio descritto da Yang et al., 1991e in cui è necessario calcolare preliminarmente il numero minimo di sezioni di conteggio in modo da intercettare ogni flusso *od* attraverso la risoluzione di un problema di programmazione lineare.

$$\begin{aligned}
\min Z(\mathbf{z}) &= \sum_l z_l \\
&\text{soggetta ai vincoli:} \\
\sum_{l \in L} \delta_{l,od} \cdot z_l &\geq 1, \quad od \in OD \\
z_l &\in [0,1]
\end{aligned} \tag{2.20}$$

dove:

z_l è una variabile binaria che assume il valore 1 se sull'arco l vi è una sezione di conteggio, 0 altrimenti;

$\delta_{l,od}$ è una variabile binaria che assume il valore 1 se vi è almeno un percorso che connette la coppia od passante per l'arco l , 0 altrimenti;

OD rappresenta l'insieme di tutte le coppie od ;

L è l'insieme di tutti gli archi presenti sulla rete;

l è il generico arco;

$Z(\mathbf{z})$ è il vettore che contiene gli identificativi degli archi a cui è associata una postazione di conteggio.

Attraverso questa procedura tutte le coppie od sono trattate allo stesso modo indipendentemente dal valore assunto dal flusso che si sposta sugli archi della rete. Chiaramente una coppia od a cui è associata una stima a priori della domanda bassa dovrebbe avere meno peso rispetto ad un'altra con flusso maggiore. Infatti, un secondo criterio che pesa le sezioni di conteggi in funzione del flusso transitante ottenuto assegnando la stima a priori della domanda è descritto da Yang e Zhou (Yang & Zhou, 1998).

$$\begin{aligned}
\max F(\mathbf{z}) &= \sum_{r \in R} f_r \cdot y_r \\
&\text{Soggetto ai seguenti vincoli:} \\
\sum_{l \in L} z_l &= \hat{l} \\
\sum_{l \in r} z_l &\geq y_r, r \in R \\
\sum_{l \in L} \delta_{l,od} \cdot z_l &\geq 1, \quad od \in OD \\
z_l \in [0,1], y_r &\in [0,1], \quad r \in R
\end{aligned} \tag{2.21}$$

dove:

f_r è il flusso di percorso associato al percorso r

y_r è una variabile binaria che assume il valore 1 quando sul percorso r è posizionata almeno una sezione di conteggio, 0 altrimenti;

$\hat{l}$ rappresenta il numero minimo di sezioni di conteggio ricavate, ad esempio, attraverso la risoluzione del problema di programmazione lineare precedente (2.21).

r e R rappresentano rispettivamente il generico percorso e l'insieme dei percorsi

2.4.3. Stimatori per la correzione della domanda

In statistica si definisce “stimatore” una funzione che associa ad ogni campione un valore del parametro da stimare. Il valore assunto dallo stimatore in corrispondenza di un particolare campione si definisce stima.

Lo stimatore principalmente utilizzato per eseguire la correzione in contesti statici del vettore di domanda $\mathbf{d}$ con i conteggi di traffico è lo stimatore dei minimi quadrati generalizzato *GLS* (*Generalized Least Squares*) (2.23). L'obiettivo di tale stimatore è identificare un nuovo vettore di domanda $\mathbf{d}^{GLS}$ che:

- 1) sia più vicino possibile alla stima a priori del vettore di domanda $\hat{\mathbf{d}}$;
- 2) assegnato alla rete generi dei flussi di arco il più vicini possibili a quelli osservati $\hat{\mathbf{f}}_c$.

$$\mathbf{d}^{GLS} = \underset{\mathbf{x} \geq 0}{argmin} \left[(\hat{\mathbf{d}} - \mathbf{x})^T \cdot \boldsymbol{\Sigma}_d^{-1} \cdot (\hat{\mathbf{d}} - \mathbf{x}) + (\hat{\mathbf{f}}_c - \mathbf{M}^c \mathbf{x})^T \cdot \boldsymbol{\Sigma}_f^{-1} \cdot (\hat{\mathbf{f}}_c - \mathbf{M}^c \mathbf{x}) \right] \quad (2.22)$$

dove:

$\hat{\mathbf{d}}$ è la stima a priori del vettore di domanda;

$\mathbf{x}$ è il vettore di domanda incognito;

$\boldsymbol{\Sigma}_d$ è la matrice di dispersione della domanda avente dimensione $n_{od} \cdot n_{od}$;

n_{od} è il numero di coppie od ;

$\boldsymbol{\Sigma}_f$ è la matrice di dispersione dei flussi d'arco avente dimensione $n_c \cdot n_c$;

n_c è il numero di archi appartenenti al sottoinsieme di sezioni di conteggio;

$\mathbf{M}^c$ è la matrice di assegnazione riferita alle sole sezioni di conteggio

$\hat{\mathbf{f}}_c$ è il vettore dei flussi d'arco osservati

La matrice di assegnazione $\mathbf{M}$ è una matrice i cui elementi m_{od}^l rappresentano l'aliquota del flusso od transitante su un determinato arco l secondo la seguente relazione:

$$m_{od}^l = \frac{f_{l,od}}{d_{od}} \quad (2.23)$$

dove:

$f_{l,od}$ rappresenta il flusso sull'arco l ottenuto assegnando il flusso di domanda, d_{od} , relativo alla generica coppia od .

Relativamente allo stimatore GLS , esiste anche una forma chiusa dello stimatore che permetti di indentificare immediatamente, attraverso l'applicazione dell'equazione di Aitken (2.25), il valore del vettore di domanda cercato senza utilizzare una procedura iterativa di ottimizzazione (Cascetta, 1984).

$$\mathbf{d}^{GLS} = (\boldsymbol{\Sigma}_d^{-1} + \mathbf{M}^{cT} \cdot \boldsymbol{\Sigma}_f^{-1} \cdot \mathbf{M}^c)^{-1} \cdot (\boldsymbol{\Sigma}_d^{-1} \cdot \hat{\mathbf{d}} + \mathbf{M}^{cT} \cdot \boldsymbol{\Sigma}_f^{-1} \cdot \hat{\mathbf{f}}_c) \quad (2.24)$$

Tale stimatore, però, non è capace di inglobare il vincolo di non negatività nelle soluzioni.

Un altro stimatore utilizzato per la correzione della domanda di trasporto è lo stimatore di massima verosimiglianza (ML) $\mathbf{d}^{ML}$ (Cascetta & Nguyen, 1988). Tale stimatore del vettore di domanda si ottiene massimizzando la probabilità di osservare i risultati delle indagini campionarie e dei flussi conteggiati sugli archi della rete. Facendo l'ipotesi che le due probabilità siano indipendenti (ipotesi piuttosto accettabile), lo stimatore di massima verosimiglianza sotto le precedenti assunzioni assume la seguente espressione:

$$\mathbf{d}^{ML} = \underset{\mathbf{x} \geq 0}{\operatorname{argmax}} [\ln L(\mathbf{n}|\mathbf{x}) + \ln L(\hat{\mathbf{f}}_c|\mathbf{x})] \quad (2.25)$$

dove:

$\mathbf{x}$ è il vettore di domanda incognito;

$\mathbf{n}$ è il vettore dei conteggi di domanda osservati;

$\hat{\mathbf{f}}_c$ è il vettore dei conteggi di traffico;

$\ln L(\mathbf{n}|\mathbf{x})$ è la funzione log-likelihood dei conteggi di domanda, cioè, il logaritmo della probabilità di osservare il vettore campionario $\mathbf{n}$ se $\mathbf{x}$ è il vettore di domanda vero;

$\ln L(\hat{\mathbf{f}}_c|\mathbf{x})$ è la funzione log-likelihood dei conteggi di traffico, cioè, il logaritmo della probabilità di osservare il vettore dei flussi d'arco $\hat{\mathbf{f}}_c$ se $\mathbf{x}$ è il vettore di domanda vero.

La distribuzione della funzione di probabilità della stima a priori della domanda dipende dalla tipologia del metodo utilizzato per stimare la domanda. Invece, i conteggi di traffico si assumono, in genere, distribuiti come variabili di Poisson oppure come variabili normali multivariate.

L'ultima tipologia di stimatori utilizzati per eseguire la correzione della domanda di trasporto sono gli stimatori bayesiani (Maher, 1983), che consentono di ottenere delle stime di parametri incogniti combinando informazioni sperimentali o campionarie con altre non sperimentali o "soggettive":

$$\mathbf{d}^b = \underset{\mathbf{x} \geq 0}{\operatorname{argmax}} [\ln g(\mathbf{x}|\hat{\mathbf{d}}) + \ln L(\hat{\mathbf{f}}_c|\mathbf{x})] \quad (2.26)$$

dove:

$\mathbf{d}^b$ rappresenta il vettore di domanda incognito ricavato attraverso lo stimatore bayesiano; $g(\mathbf{x}|\hat{\mathbf{d}})$ rappresenta la funzione di probabilità a priori della domanda che può essere specificata in diversi modi. Ad esempio, si può assumere che il vettore di domanda incognito sia distribuito come una variabile aleatoria normale multivariata;

$L(\hat{\mathbf{f}}_c|\mathbf{x})$ rappresenta la funzione di probabilità a priori sui flussi d'arco.

Relativamente al problema di correzione del vettore di domanda $\mathbf{d}$, le informazioni sperimentali sono reperibili attraverso i conteggi di traffico, mentre le informazioni non sperimentali possono essere, ad esempio, una vecchia stima della domanda, oppure una stima della domanda ottenuta da modello.

Gli stimatori bayesiani sono ottenuti a partire dalla funzione di probabilità a posteriori del vettore di domanda incognito $\mathbf{x}$, condizionato alle informazioni a priori e quelle sperimentali.

I tre stimatori sotto certe ipotesi possono essere ritenuti equivalenti, ad esempio, se si ipotizza che i conteggi di domanda $\mathbf{n}$ e di traffico $\hat{\mathbf{f}}_c$ dell'equazione (2.25) siano delle variabili aleatorie distribuite come Normali Multivariate, con medie rispettivamente pari a $\hat{\mathbf{d}}$ e $\mathbf{M}^c \cdot \mathbf{x}$, matrici di covarianza pari a Σ_f e Σ_d , lo stimatore di massima verosimiglianza è equivalente allo stimatore GLS. Anche nel caso di stimatori bayesiani, se le funzioni di probabilità a priori della domanda $g(\mathbf{x}|\hat{\mathbf{d}})$ e dei flussi d'arco $L(\hat{\mathbf{f}}_c|\mathbf{x})$ dell'equazione (2.26) sono distribuite come Normali Multivariate si ottiene un problema formalmente analogo a quello definito attraverso lo stimatore GLS.

Dunque, gli stimatori bayesiani e di massima verosimiglianza sono riconducibili allo stimatore GLS, il quale ha un notevole vantaggio rispetto ai precedenti, cioè, non necessita di alcuna ipotesi sulla distribuzione delle stime. Inoltre, sulla base di numerosi test si è dimostrato funzionare meglio rispetto agli altri.

3. Correzione della matrice di domanda mediante clustering sequenziale di flussi *od*

L'obiettivo del presente capitolo è quello di presentare una metodologia di correzione della matrice *od* mediante conteggi di traffico. Le performance della procedura di correzione sono influenzate dal rapporto r tra numero di equazioni n_c e il numero incognite n_{od} .

Nei casi reali il numero di incognite è di gran lunga superiore al numero di equazioni linearmente indipendenti, rendendo il sistema fortemente indeterminato, poiché caratterizzato da $\infty^{(n_{od}-n_c)}$ soluzioni.

Studi specifici di letteratura (Simonelli et al., 2012), hanno mostrato chiaramente con esperimenti di laboratorio, che la procedura di correzione della matrice *od* funziona in modo efficace solamente quando:

- il rapporto r tra le equazioni linearmente indipendenti e le incognite del problema è prossimo all'unità: al diminuire di tale rapporto, invece, la procedura tende a generare una matrice *od* sempre più vicina alla stima iniziale – preservandone quindi gli errori – modificata solo del minimo necessario per riprodurre i conteggi di traffico disponibili.
- il posizionamento dei conteggi di traffico garantisce un buon livello di coverage per tutte le coppie *od*: infatti, se una coppia *od* non è intercettata da nessuna sezione di conteggio, non può essere corretta.

Sulla base di queste premesse, il presente capitolo di tesi si propone l'obiettivo di individuare approcci innovativi di correzione della matrice *od*, che consentano a tale procedura di “lavorare” nelle condizioni ottimali appena descritte.

3.1. Descrizione della metodologia

In questo paragrafo è illustrato l'approccio utilizzato per migliorare l'efficacia dello stimatore *GLS* nella procedura di correzione della domanda con conteggi di traffico. La metodologia proposta si basa su un algoritmo euristico sequenziale in grado di eseguire dei clustering sequenziali mediante aggregazione di coppie *od* e dei relativi flussi *od*. Gli input iniziali dell'algoritmo sono: una stima a priori del vettore di domanda $\mathbf{d}^0$ con cardinalità n_{od} coerente con un set iniziale di coppie *od* P^0 , un vettore archi conteggiati $\mathbf{C}$, una matrice di assegnazione $\mathbf{M}^{C0}$, un numero finale di coppie *od* n_{od}^* da raggiungere. La struttura dell'algoritmo è la seguente, dove con i si indica la sequenza dei passi di aggregazione:

Step 1: Aggregazione

Per $i=1$ a $n_{od}-n_{od}^*$

- a) Selezione delle p_1 e p_2 coppie *od* da aggregare, con $p_1, p_2 \in P^{i-1}$. Vale la pena sottolineare che l'algoritmo lavora all'iterazione i con un insieme di coppie *od* già aggregate, denotato con P^{i-1} a cardinalità $|P^{i-1}| = n_{od}-i+1$;
- b) Aggregazione dei flussi *od* e della matrice di assegnazione coerentemente con l'aggregazione di p_1 e p_2 , ottenendo rispettivamente $\mathbf{d}^i$ da $\mathbf{d}^{i-1}$ e $\mathbf{M}^{Ci}$ da $\mathbf{M}^{Ci-1}$.
- c) Aggregazione della matrice di dispersione della domanda coerentemente con l'aggregazione di p_1 e p_2 .

Step 2: Correzione dei flussi *od* aggregati

Correzione del vettore dei flussi *od* aggregati $\mathbf{d}^{n_{od}-n_{od}^*}$ con cardinalità $n_{od}^* = |P^{n_{od}-n_{od}^*}|$ utilizzando i conteggi $\mathbf{c}$, la matrice di assegnazione $\mathbf{M}^{C^{n_{od}-n_{od}^*}}$ e lo stimatore *GLS* in modo da ottenere un vettore aggregato e corretto $\mathbf{d}^{GLS, n_{od}-n_{od}^*}$.

Step 3: Disaggregazione del flusso *od* corretto

Disaggregazione del vettore di domanda corretto $\mathbf{d}^{GLS, n_{od}-n_{od}^*}$ fino alla granularità iniziale costituita da P^0 coppie *od*, ottenendo così $\mathbf{d}^{GLS, 0}$.

Con riferimento allo step 1a, l'aggregazione di due coppie od (p_1 e p_2) all' i -esima iterazione si basa su un doppio criterio, uno di adiacenza topologica e un altro di coverage. La coverage χ_{od}^i per la generica coppia od all'iterazione i -esima è definita come nella 2.17:

$$\chi_{od}^i = \sum_{c \in \mathcal{C}} m_{od}^c{}^{i-1} \quad m_{od}^c{}^{i-1} \in \mathbf{M}^{C_{i-1}} \quad (3.1)$$

cioè come somma degli elementi della colonna di $\mathbf{M}^C$ relativa a quella specifica coppia od , ed è quindi maggiore di zero solo se il flusso relativo a quella coppia od è intercettato in almeno una sezione di conteggio.

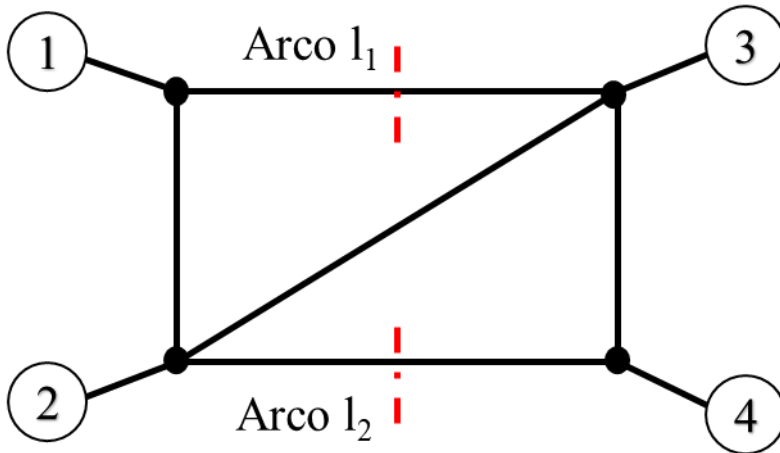
La metodologia di aggregazione utilizzata mira ad aggregare le coppie od con copertura nulla o bassa con coppie od ad alta copertura. Per fare ciò, la (3.1) può essere applicata direttamente a tutte le coppie od appartenenti a P^{i-1} , prima ordinandole in senso decrescente in base alla copertura, in modo che l'ultima coppia od sia definita come $p_{last} = \{o_{last}, d_{last}\} \in P^{i-1}$. Si potrebbe anche verificare una situazione con più coppie od a copertura nulla, in questi casi, p_{last} può essere selezionato o casualmente oppure scegliendo la coppia od con copertura nulla adiacente alla coppia od con copertura massima. Seguendo un criterio di adiacenza, l'insieme P_{adj} delle coppie od adiacenti a p_{last} rappresenta l'insieme delle coppie od candidate all'aggregazione con p_{last} , cioè P_{adj} include tutte le coppie od $p \equiv \{o, d\}$ tali che $o \in A_{o_{last}}$ e $d \in A_{d_{last}}$, dove $A_{o_{last}}$ e $A_{d_{last}}$ rappresentano gli insiemi contenenti le zone adiacenti rispettivamente alla zona o_{last} e alla zona d_{last} .

Infine, p_{last} è aggregato alla coppia od $p^* \in P_{adj}$ con copertura massima. Anche in questo caso, p^* potrebbe non essere unico, e sarebbe sufficiente una selezione casuale.

In particolare, anche se non necessario, il ricorso a un criterio di adiacenza ha l'obiettivo primario di evitare un eccessivo carico computazionale. A questo proposito, vale la pena notare che l'approccio proposto produce un'aggregazione di coppie od senza alcuna interpretazione topologica/fisica: piuttosto, dovrebbe essere visto solo come un veicolo matematico per aumentare l'efficacia della procedura di correzione della domanda prima di riportare le copie od alla granularità iniziale.

In figura 3.1 è riportato un esempio di clusterizzazione applicata ad una rete molto semplice costituita da quattro coppie od , nove archi bidirezionali e due sezioni di conteggio. Nell'esempio, al primo cluster sono state associate le coppie od (1-3) e (1-4), mentre nel

secondo (2-3) e (2-4) coerentemente con la metodologia spiegata precedentemente nel caso più generale.



Zone adiacenti:

1 e 2

3 e 4

M	od_{1-3}	od_{2-3}	od_{1-4}	od_{2-4}
Arco l_1	0,7	0,2	0,1	0,2
Arco l_2	0,2	0,2	0,6	0,7

Cluster 1

(od_{1-3}, od_{1-4})

Cluster 2

(od_{2-3}, od_{3-4})

Fig. 3.1 Esempio di clusterizzazione su rete bidirezionale con 2 sezioni di conteggio

Una volta selezionate le coppie od p_{last} e p^* da aggregare, sia il vettore dei flussi od $\mathbf{d}^{i-1}$ che la corrispondente matrice di assegnazione $\mathbf{M}^{c_{i-1}}$ devono essere modificati coerentemente. Per comodità di notazione, denotiamo p_{last} e p^* rispettivamente con p_1 , $p_2 \in P^{i-1}$ e identifichiamo con $p_{12} = \{o_{12}, d_{12}\} \in P^i$ la coppia od aggregata risultante. La correzione del vettore od richiede semplicemente la somma dei flussi od relativi a p_1 e p_2 , ottenendo un nuovo vettore di flusso od $\mathbf{d}^i$ tale che la dimensione di $\mathbf{d}^i$ sia uguale alla cardinalità di P^i e la dimensione di $\mathbf{d}^{i-1}$ alla cardinalità di P^{i-1} . L'aggregazione della matrice di assegnazione richiede l'aggregazione di tutti i termini $m_{od}^{c_{i-1}} \in \mathbf{M}^{c_{i-1}}$ con $od = \{o_1 d_1, o_2 d_2\} \forall c$. Ciò può essere ottenuto facilmente attraverso una media ponderata di tutte le coppie interessate $m_{o_1 d_1}^{c_{i-1}}$ e $m_{o_2 d_2}^{c_{i-1}} \forall c$ basata sui corrispondenti flussi od a priori $d_{o_1 d_1}^{prior, i-1}$ e $d_{o_2 d_2}^{prior, i-1}$, ottenendo i nuovi termini $m_{o_{12} d_{12}}^c \forall c$:

$$\begin{aligned}
m_{o_{12}d_{12}}^c{}^i &= m_{o_1d_1}^c{}^{i-1} \frac{d_{o_1d_1}^{prior,i-1}}{d_{o_1d_1}^{prior,i-1} + d_{o_2d_2}^{prior,i-1}} \\
&+ m_{o_2d_2}^c{}^{i-1} \frac{d_{o_2d_2}^{prior,i-1}}{d_{o_1d_1}^{prior,i-1} + d_{o_2d_2}^{prior,i-1}}
\end{aligned} \tag{3.2}$$

L'equazione precedente introduce chiaramente un errore, perché le proporzioni di flusso od sono aggregate usando la stima a priori della domanda, distorta per definizione. Tuttavia, la controparte positiva per questa distorsione è che l'aggregazione potrebbe migliorare l'efficacia della correzione della domanda mediante l'applicazione dello stimatore GLS . A valle di un'aggregazione, prima di applicare la procedura di correzione, anche la matrice di dispersione della domanda deve essere aggregata, sommando le varianze delle domande associata alle coppie od aggregate.

A valle della correzione della domanda aggregata, i flussi od aggregati e corretti $\mathbf{d}^{GLS, n_{od}-n_{od}^*}$ sono disaggregati alla granularità iniziale fino a $\mathbf{d}^{GLS, 0}$, cioè alla granularità dell'insieme iniziale di coppie od P^0 , conservando le proporzioni relative calcolate sulla stima a priori del vettore di domanda. In particolare, sia $d_{o_{agg}d_{agg}}^{GLS} \in \mathbf{d}^{GLS, n_{od}-n_{od}^*}$ il generico flusso od aggregato e corretto e $P_{o_{agg}d_{agg}} \subseteq P^0$ un insieme delle coppie od iniziali incluse in $\{o_{agg}, d_{agg}\}$; si ha:

$$\begin{aligned}
d_{od}^{GLS} &= d_{o_{agg}d_{agg}}^{GLS} \cdot \frac{d_{od}^{prior}}{\sum_{o'd' \in P_{o_{agg}d_{agg}}} d_{o'd'}^{prior}} \quad \forall \{o, d\} \\
&\in P_{o_{agg}d_{agg}} \quad \forall d_{o_{agg}d_{agg}}^{GLS} \in \mathbf{d}_{GLS}^{n_{od}-n_{od}^*}
\end{aligned} \tag{3.3}$$

Nel complesso, l'euristica proposta permette di bilanciare efficacemente incognite ed equazioni nella procedura di correzione dei flussi od al prezzo di introdurre un errore nella procedura di assegnazione. Come già menzionato nell'introduzione, l'obiettivo di questa metodologia è di esplorare la fattibilità del trade-off tra un migliore correzione dei flussi od e la distorsione introdotta dalle fasi di aggregazione-disaggregazione.

Di seguito è riportato lo schema della procedura proposta.

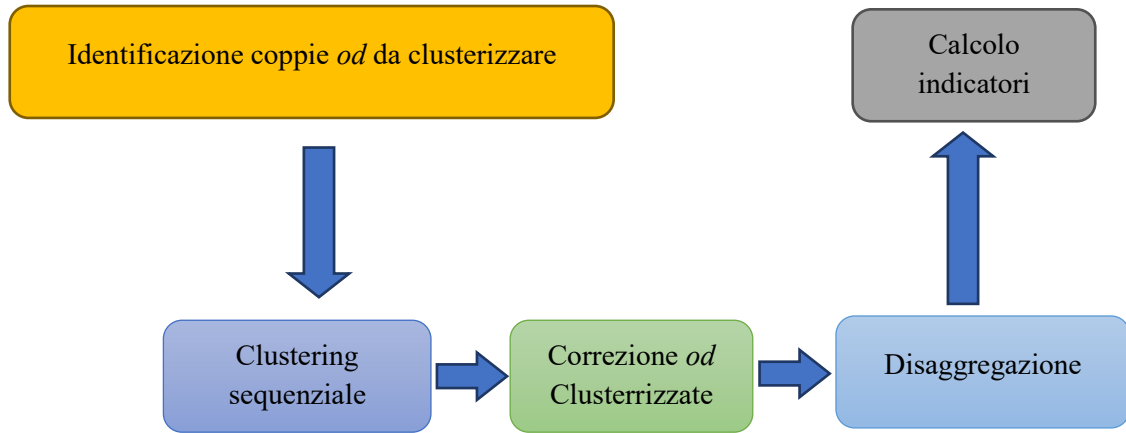


Fig. 3.2 Schema della metodologia

A valle di ogni disaggregazione sono stati calcolati una serie di indicatori prestazionali, in particolare, lo stimatore utilizzato per il calcolo degli indicatori è il *cvRMSE* definito nel seguente modo:

$$cvRMSE = \frac{\sqrt{\frac{\sum_{i=1}^N (X_i^{True} - X_i^{Update})^2}{N}}}{\sum_{i=1}^N X_i^{True} / N} \quad (3.4)$$

dove:

X_i^{True} è la generica variabile (domanda, flussi d'arco, etc.) utilizzata come base di confronto;

X_i^{Update} è la generica variabile ottenuta a valle della procedura di correzione;

N rappresenta il numero di variabili considerate (numero di coppie *od*, numero di archi, etc.)

3.2. Descrizione del caso studio

La procedura precedentemente descritta è stata applicata in prima analisi ad una toy network e successivamente ad una rete reale rappresentata dal comune di Caserta.

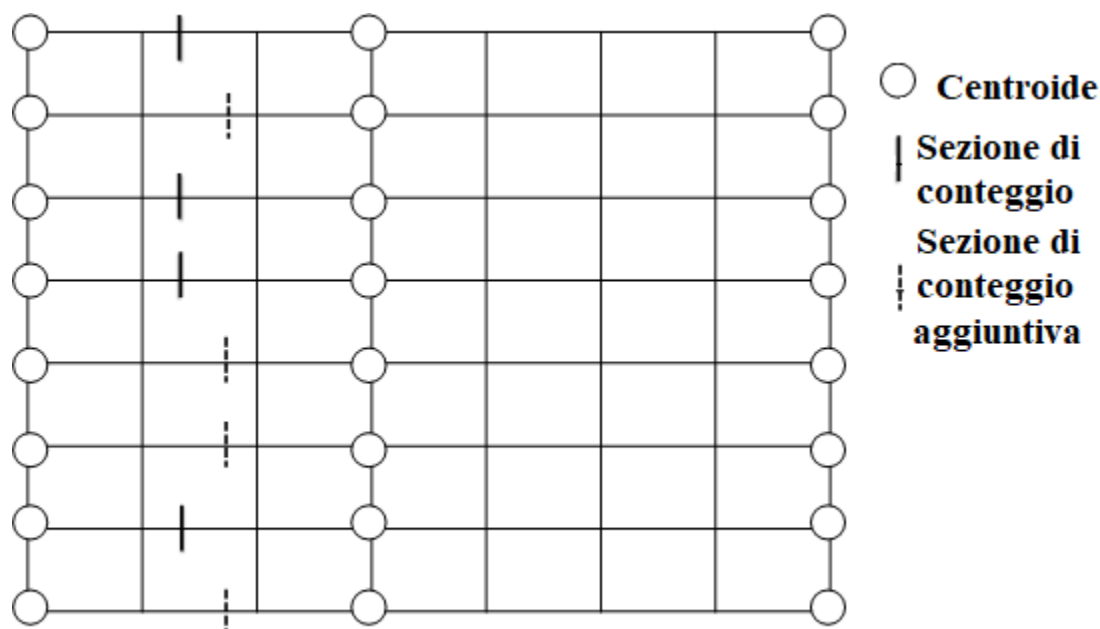


Fig. 3.3 Rete test

La toy network a cui è stata applicata la metodologia descritta nel paragrafo precedente è costituita da una griglia di dimensione 7x7 con archi bidirezionali. Inoltre, vi sono 24 centroidi e 8 sezioni di conteggio (Fig. 3.3). Nel complesso, sono stati presi in considerazione solo i flussi *od* diretti da sinistra verso destra, ottenendo 192 flussi *od* non nulli, cioè 128 flussi *od* in uscita dai centroidi del lato di sinistra (64 flussi *od* ciascuno rispettivamente verso la colonna di centroidi centrale e verso la colonna di centroidi di destra) e 64 flussi *od* dalla colonna dei centroidi centrale verso la colonna dei centroidi di destra. In particolare, in un primo esperimento sono state utilizzate solo 4 sezioni di conteggio (linea continua), mentre in un secondo esperimento sono state considerate le ulteriori 4 sezioni di conteggio (linea tratteggiata) per un totale di 8 sezioni di conteggio, simulando in questo secondo caso una screen-line di sezioni di conteggio.

Per l'applicazione della procedura si sono ipotizzati i seguenti settaggi sui parametri:

La domanda di trasporto ipotizzata vera d^{True} è stata ottenuta estraendo dei flussi *od* da una distribuzione di tipo normale a media 500 e coefficiente di variazione pari a 0.05.

La stima a priori del vettore di domanda $\mathbf{d}^{Prior}$ è stata ottenuta perturbando la domanda ipotizzata vera, inoltre, per una maggiore generalità dei risultati, sono state generate $p \in 1 \dots n_p$ replicazioni della stima a priori utilizzando due diverse distribuzioni una uniforme con media d_p^{Mean} definita nell'intervallo $d_p^{Mean} \pm d_p^{Mean} \cdot k_p$, a cui corrisponde una varianza della singola coppia od paria a

$$\sigma_{od,p}^2 = \frac{(2d_p^{Mean} \cdot k_p)^2}{12} = \frac{(d_p^{Mean} \cdot k_p)^2}{3} \quad (3.5)$$

ed una normale:

$$\begin{aligned} \mathbf{d}_p &\sim N(\mathbf{d}_p^{Mean}, \mathbf{\Sigma}_{d,p}) \\ \sigma_{od,p}^2 &= \frac{(d_p^{Mean} \cdot k_p)^2}{16} \\ \mathbf{d}_p^{Mean} &= \mathbf{d}^{True} \cdot (1 + k_p \cdot \mathbf{v}) \\ k_p &\sim U[0,1] \\ \mathbf{v} &\sim U[0,1] \end{aligned} \quad (3.6)$$

dove:

$\mathbf{d}_p$ rappresenta la stima a priori del vettore di domanda alla generica replicazione p ;

d_p^{Mean} rappresenta la media della stima a priori nella replicazione p ;

$\mathbf{\Sigma}_{d,p}$ rappresenta la matrice di dispersione della stima della domanda alla replicazione p ;

k_p è un numero estratto casualmente da una distribuzione di tipo uniforme nell'intervallo $[0,1]$, costante per la replicazione p ;

$\mathbf{v}$ è un vettore di dimensione $n_{od} \cdot 1$, costante per ogni replicazione, estratto da una distribuzione di tipo uniforme nell'intervallo $[0,1]$.

n_p rappresenta il numero di repliche fissato a 100.

Le scelte di varianza effettuate, in entrambi i casi, hanno l'obiettivo di evitare valori negativi di domanda.

La matrice di assegnazione $\mathbf{M}$ è stata ottenuta attraverso una procedura di carico stocastico della rete (SNL) a doppio passo con modello di scelta del percorso di tipo *logit*. Inoltre, la matrice di assegnazione è stata ipotizzata “error-free” cioè stima della matrice di assegnazione utilizzata in fase di correzione $\hat{\mathbf{M}}$ coincidente con la matrice di assegnazione ipotizzata vera $\mathbf{M}$.

Lo stimatore della domanda utilizzato è il *GLS* applicato mediante la sua formulazione in forma chiusa (3.6).

$$\mathbf{d}_p^{GLS} = (\boldsymbol{\Sigma}_{d,p}^{-1} + \mathbf{M}^{c^T} \cdot \boldsymbol{\Sigma}_f^{-1} \cdot \mathbf{M}^c)^{-1} \cdot (\boldsymbol{\Sigma}_{d,p}^{-1} \cdot \mathbf{d}_p + \mathbf{M}^{c^T} \cdot \boldsymbol{\Sigma}_f^{-1} \cdot \hat{\mathbf{f}}_c) \quad (3.7)$$

La matrice di dispersione della domanda $\boldsymbol{\Sigma}_{d,p}$, in mancanza di una misura relativa alla “fiducia” nella stima a priori della domanda, è stata ipotizzata coincidente con la matrice identità. Invece, la matrice di dispersione dei flussi di arco $\boldsymbol{\Sigma}_f$ è stata assunta diagonale ed eteroschedastica con varianza proporzionale ai flussi d’arco conteggiati. Per definire tale coefficiente di proporzionalità, con il principale obiettivo di tenere conto del diverso livello di affidabilità delle informazioni a priori, si è effettuata un’analisi parametrica che ha consentito di stabilire che $\boldsymbol{\Sigma}_{d,p}$ è il caso sia più grande di due ordini di grandezza circa rispetto a $\boldsymbol{\Sigma}_f$.

La seconda rete a cui è stata applicata la procedura è quella di Caserta, costituita da 65 centroidi, circa 9000 archi e due set di conteggi. Il primo set di conteggi è costituito da 25 archi, mentre il secondo set da 18 archi in aggiunta ai precedenti per un totale di 43 archi.

I settaggi utilizzati in questo caso sono gli stessi utilizzati per la toy network, tranne per la media della distribuzione dalla quale è stato estratto il vettore di domanda vero $\mathbf{d}^{True}$ fissata a 5 (invece di 500) e il numero di replicazioni n_p fissate a 10 (invece di 100).

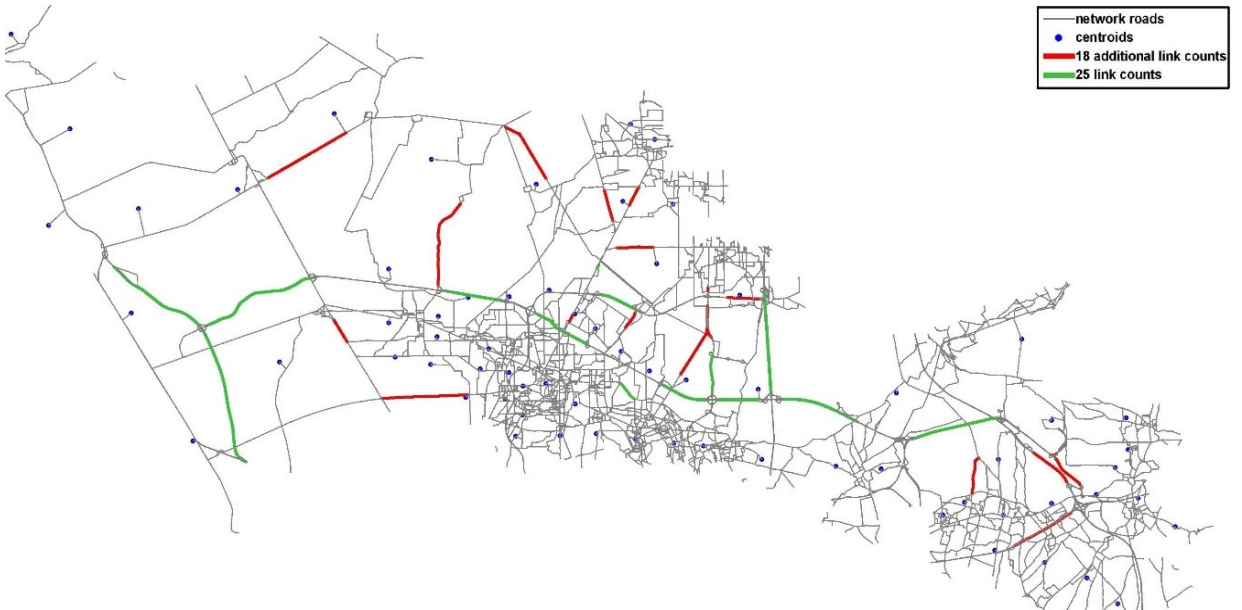


Fig. 3.4 Rete test di Caserta

3.3. Risultati e conclusioni

In questa sezione sono riportati i risultati degli esperimenti relativi sia alla toy-network sia alla rete reale di Caserta. L'indicatore utilizzato per definire una misura dell'errore è il $cvRMSE$ (3.4), il quale è stato valutato tra:

- la domanda ipotizzata vera e la corrispondente stima a priori aggregata-corretta-disaggregata;
- TDV^4 (total demand volume) della domanda ipotizzata vera e il TDV della stima a priori aggregata-corretta-disaggregata;
- L'hold-out dei flussi d'arco ottenuti assegnando la domanda ipotizzata vera e la domanda aggregata-corretta-disaggregata. Tale hold-out rappresenta il sottoinsieme degli archi non appartenenti alle sezioni di conteggio.
- la matrice di assegnazione aggregata sulla base della domanda ipotizzata vera e la corrispondente matrice aggregata sulla base della stima a priori della domanda.

Siccome l'intera procedura proposta è stata replicata p volte, per una migliore rappresentazione grafica sono stati riportati dei risultati sintetici in corrispondenza di tre percentili 25°, 50° e 75°, la cui descrizione da questo punto in poi avverrà rispetto al 50° percentile in quanto i tre percentili hanno andamento analogo.

I risultati sono presentati in figura 3.4 e figura 3.5 per la perturbazione della stima a priori della domanda di tipo uniforme e rispettivamente per 4 e 8 sezioni di conteggio. Analogamente in figura 3.7 e figura 3.8 per la perturbazione della stima a priori della domanda di tipo normale. L'asse delle ascisse riporta il livello di aggregazione normalizzato, esso è definito come il rapporto tra il numero di aggregazioni eseguite e il numero massimo di aggregazioni possibili. In questa rappresentazione al punto iniziale compete aggregazione normalizzata pari a 0 (nessuna aggregazione), mentre al punto finale compete un valore di aggregazione normalizzata pari a 1 (cioè una sola coppia *od* rimasta). Inoltre, tutti i diagrammi illustrano nel punto 0 il $cvRMSE$ relativo alla perturbazione iniziale (cioè senza correzione): questo permette di capire l'effetto migliorativo generato dalla correzione ai diversi livelli di aggregazione.

⁴ TDV rappresenta la domanda complessiva movimentata nel sistema. Si ottiene sommando gli elementi della matrice di domanda.

Un primo risultato degno di nota è che, qualunque sia la perturbazione iniziale, l'aggregazione migliora l'efficacia della procedura di correzione dei flussi *od*, nel caso di 4 sezioni di conteggio si osserva con una riduzione percentuale del sopra definito *cvRMSE* sui flussi *od* del 22,6% (da 0,31 senza aggregazione a 0,24 all'aggregazione massima) nel caso della perturbazione uniforme e del 23,8% (da 0,21 senza aggregazione a 0,16 all'aggregazione massima) nel caso della perturbazione normale e 4 sezioni di conteggio. Anche relativamente all'hold out dei flussi di arco, si osservano significativi miglioramenti con una riduzione percentuale dell'indicatore *cvRMSE* del 50% (da 0.12 senza aggregazione a 0.06 all'aggregazione massima) nel caso di perturbazione normale e del 54% (da 0.11 senza aggregazione a 0.05 all'aggregazione massima) nel caso di perturbazione uniforme. L'esperimento con 8 sezioni di conteggio dà risultati molto simili.

È interessante notare che l'approccio proposto è in grado di stimare molto bene il *TDV* in tutti gli esperimenti, con una riduzione percentuale del *cvRMSE* del 90.6% (da 0.08 senza aggregazione a 0.008 con aggregazione massima) nella perturbazione uniforme e del 91.3% (da 0.08 senza aggregazione a 0.007 con aggregazione massima) nel caso di perturbazione normale.

Infine, anche l'andamento della matrice di assegnazione è molto interessante. In entrambi i casi di perturbazione uniforme e normale, il *cvRMSE* della matrice di assegnazione (che parte da zero per definizione) aumenta nei passi iniziali dell'aggregazione come conseguenza attesa dell'errore introdotto dall'equazione (3.2), ma tale errore svanisce quando ci si avvicina all'aggregazione massima. Questa performance generale ha andamenti leggermente differenti negli esperimenti con 4 e 8 sezioni di conteggio rispettivamente.

Chiaramente, questi risultati sono interconnessi: la caratteristica chiave dell'approccio proposto sta nella sua capacità di stimare il *TDV* ground-truth all'aggregazione massima, molto meglio di quanto riesca a fare la procedura di correzione applicata ai flussi *od* non aggregati, a causa della sua capacità di ottenere, all'aggregazione massima, i valori aggregati corretti degli elementi della matrice di assegnazione relativi alle sezioni di conteggio. A sua volta, la stima più efficace del *TDV* permette, nella post-disaggregazione (step 3), una correzione più precisa dei flussi *od* disaggregati, con conseguente migliore capacità di riprodurre i flussi d'arco (anche sull'hold-out). Chiaramente, più grande è l'errore nel

volume di domanda totale della stima precedente, migliore è la performance dell'approccio proposto.

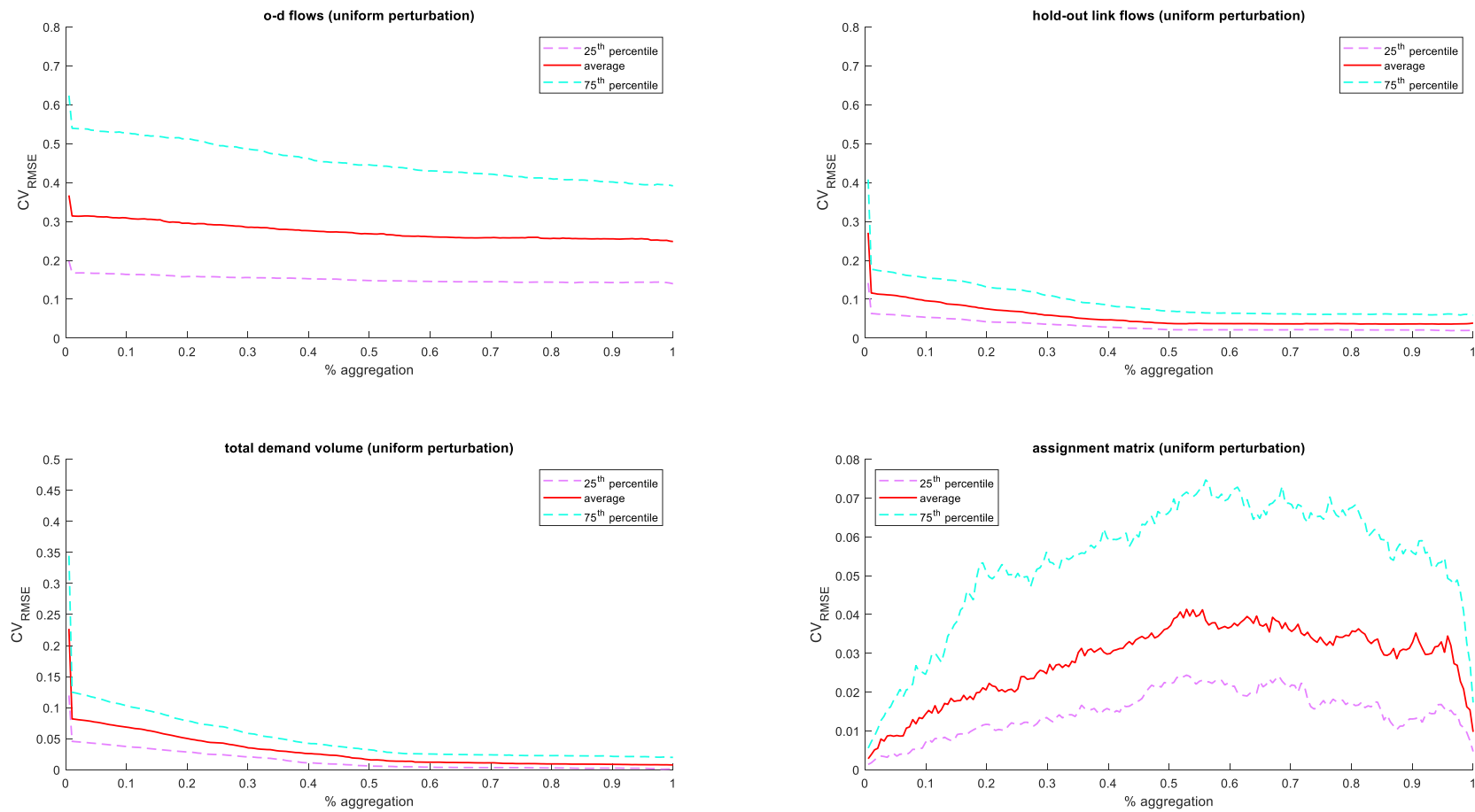


Fig. 3.5 Toy network: risultati sperimentali con perturbazione della stima a priori di tipo uniforme e 4 sezioni di conteggio

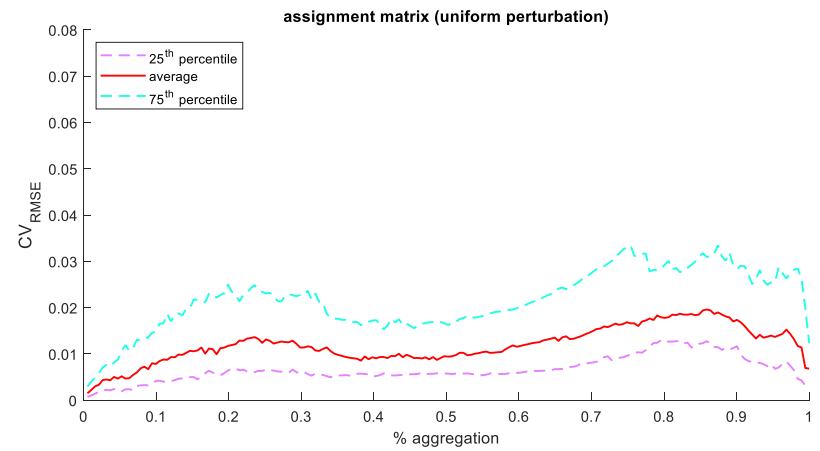
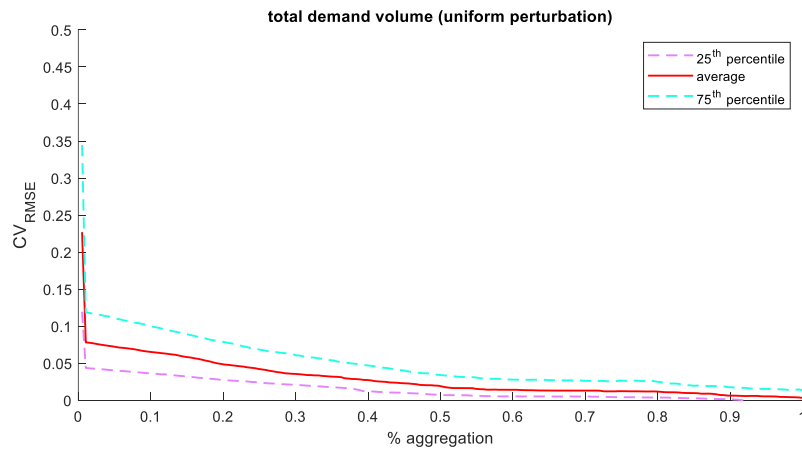
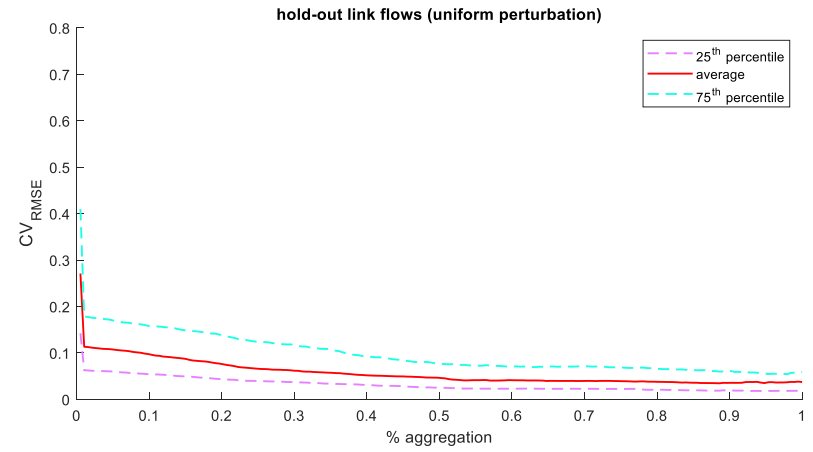
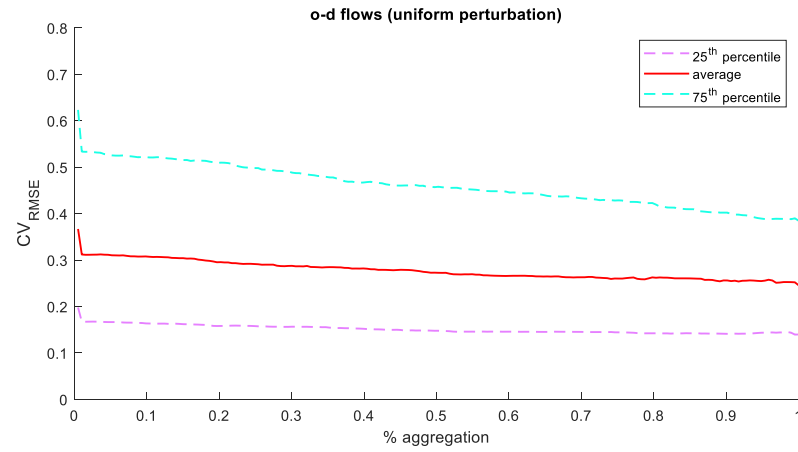


Fig. 3.6 Toy network: risultati sperimentali con perturbazione della stima a priori di tipo uniforme e 8 sezioni di conteggio

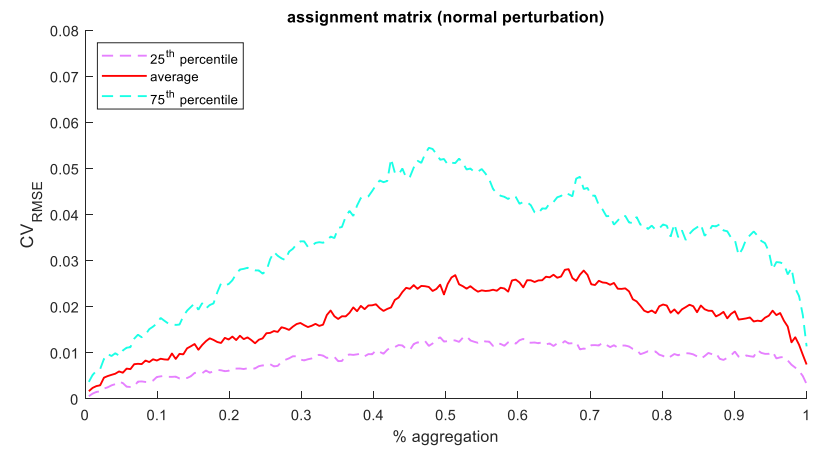
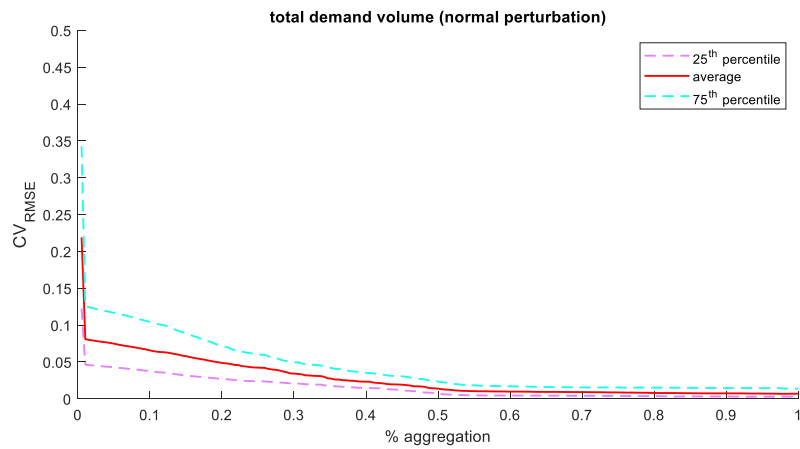
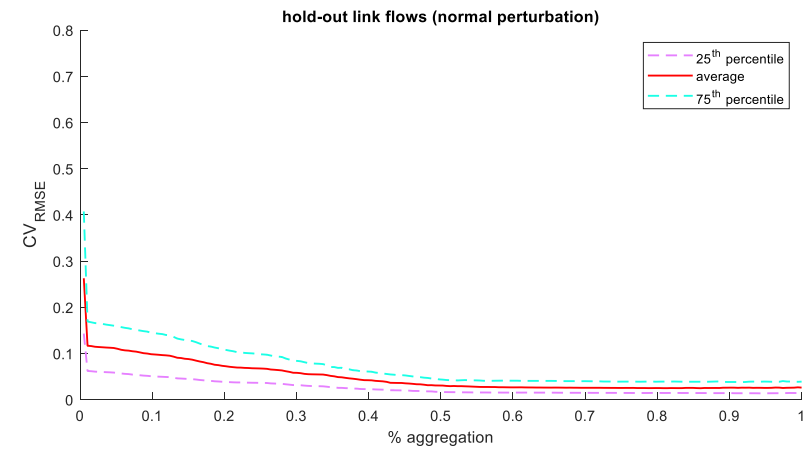
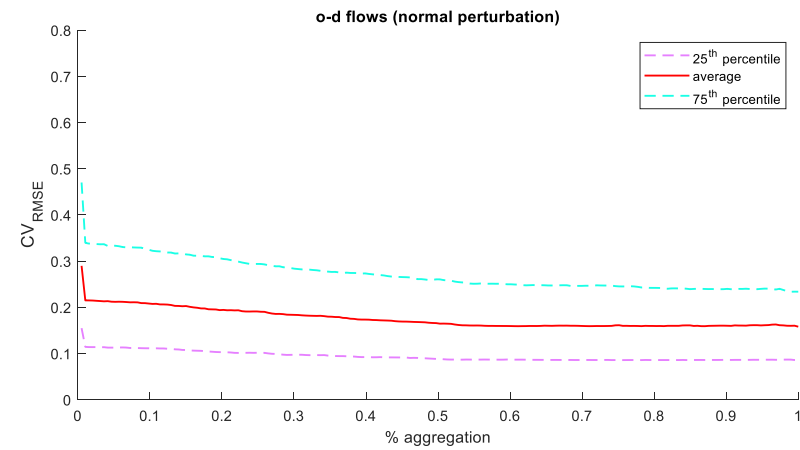


Fig. 3.7 Toy network: risultati sperimentali con perturbazione della stima a priori di tipo normale e 4 sezioni di conteggio

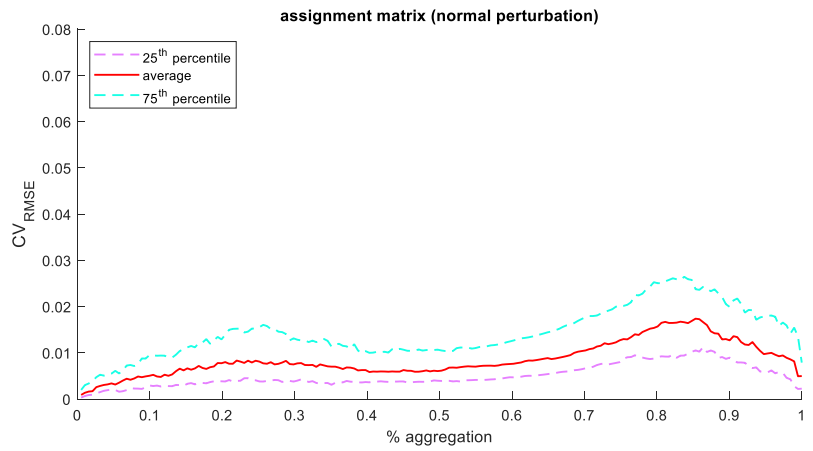
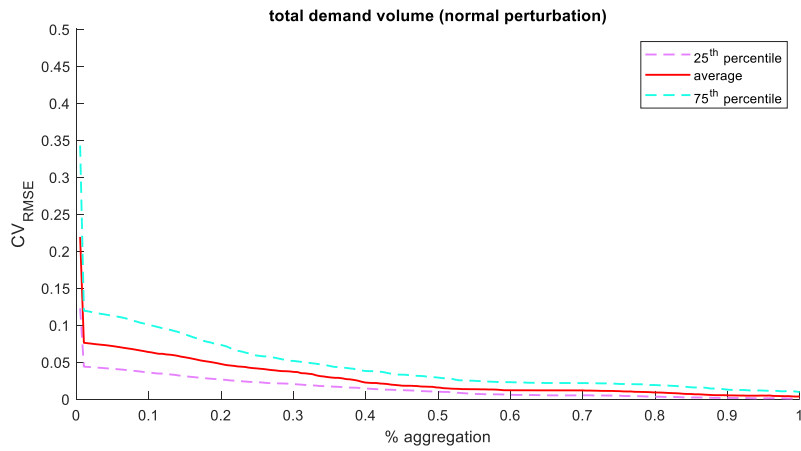
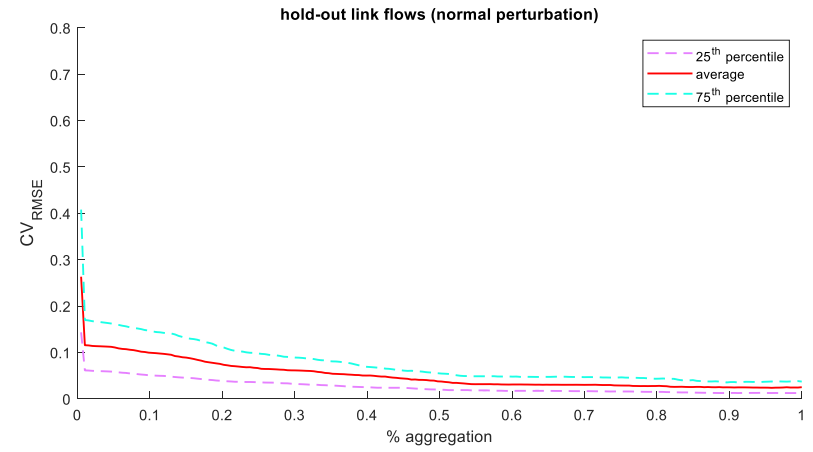
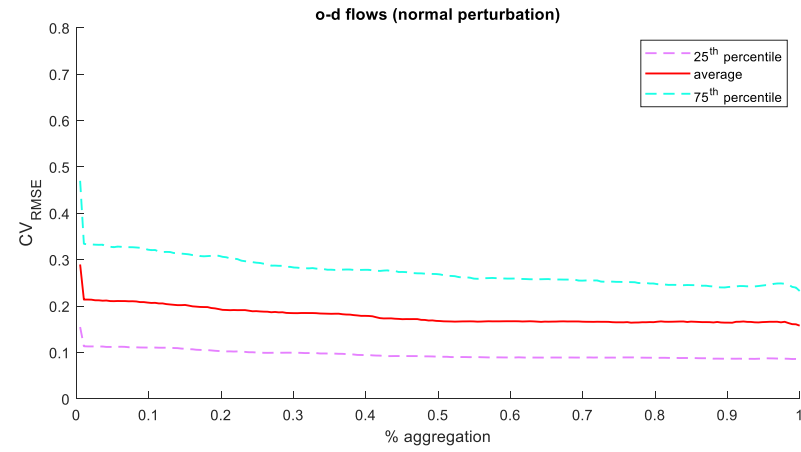


Fig. 3.8 Toy network: risultati sperimentali con perturbazione della stima a priori di tipo uniforme e 4 sezioni di conteggio

I risultati relativi all'applicazione alla rete di Caserta sono riportati in figura 3.9 e figura 3.10 per il caso di perturbazione a priori uniforme e rispettivamente 25 e 43 sezioni di conteggio, e simmetricamente in figura 3.11 e figura 3.12 per il caso di distribuzione normale.

Anche in questo caso si osserva come l'aggregazione migliori l'efficacia della procedura di correzione dei flussi *od*, con una riduzione percentuale del *cvRMSE* nel caso di distribuzione uniforme del 22.5% (da 0.4 senza aggregazione a 0.31 all'aggregazione massima) nel caso di 25 sezioni di conteggio e del 39.3% (da 0.28 senza aggregazione a 0.17 all'aggregazione massima) nel caso di 43 sezioni di conteggio. Nel caso di distribuzione normale si ha una riduzione del *cvRMSE* sui flussi *od* del 29.6% (da 0.27 senza aggregazione a 0.19 all'aggregazione massima) nel caso di 25 sezioni di conteggio e del 33.3% (da 0.27 senza aggregazione a 0.18 all'aggregazione massima) nel caso di 43 sezioni di conteggio.

È interessante notare come anche in questo caso l'approccio proposto sia in grado di stimare con buona approssimazione il *TDV* in tutti gli esperimenti, con una riduzione percentuale del *cvRMSE* del 70,0% (da 0,13 senza aggregazione a 0,039 con aggregazione massima) nel caso di 25 sezioni di conteggio e dell'80,8% (da 0,12 senza aggregazione a 0,023 con aggregazione massima) nel caso di 43 sezioni di conteggio, con perturbazione della domanda di tipo normale. I grafici relativi alla perturbazione uniforme non sono stati commentati in quanto hanno andamenti molto simili a quelli appena visti.

In definitiva l'approccio proposto ha mostrato la capacità di migliorare l'affidabilità della procedura di correzione della domanda di trasporto, in quei contesti in cui sono già disponibili delle sezioni di conteggi la cui posizione è nota a priori e quindi non modificabile. I precedenti esperimenti di laboratorio hanno evidenziato il miglioramento marginale della procedura di correzione della matrice *od* mediante una significativa riduzione dell'indicatore d'errore *cvRMSE*. La procedura testata si basa su un'euristica iterativa che in contesti reali ha chiaramente dei tempi di calcolo troppo alti per poter essere utilizzata; per questo motivo, nel capitolo successivo si proporranno metodologie di correzione sempre basate sull'aggregazione, ma implementabili con tempi di calcolo ragionevoli anche su reti reali.

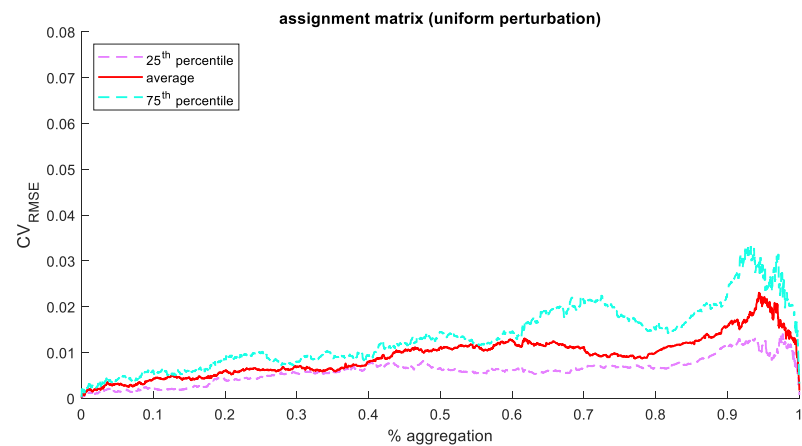
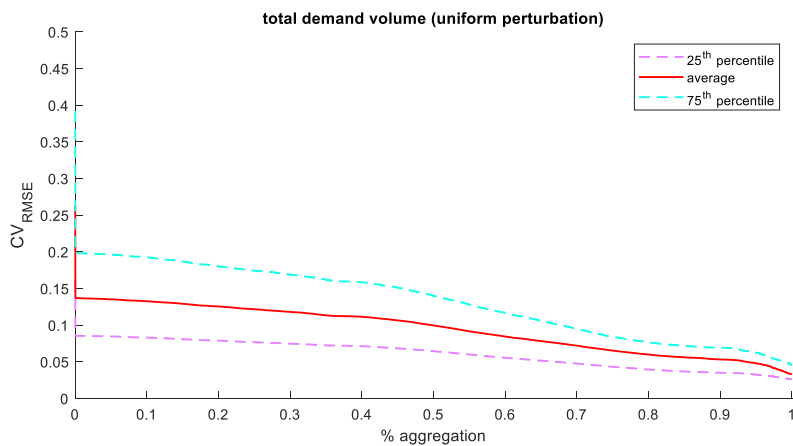
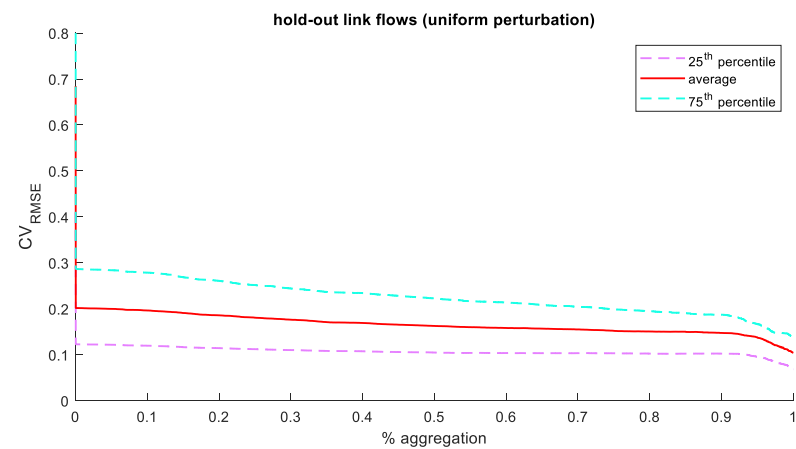
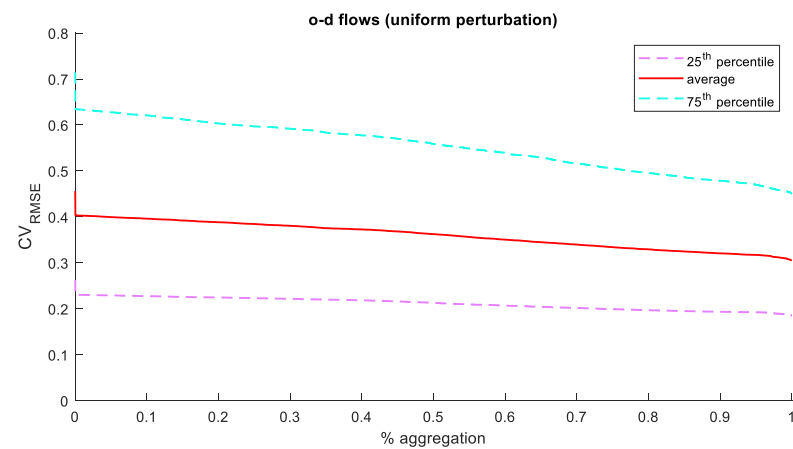


Fig. 3.9 Rete di Caserta: risultati sperimentali con perturbazione della stima a priori di tipo uniforme e 25 sezioni di conteggio

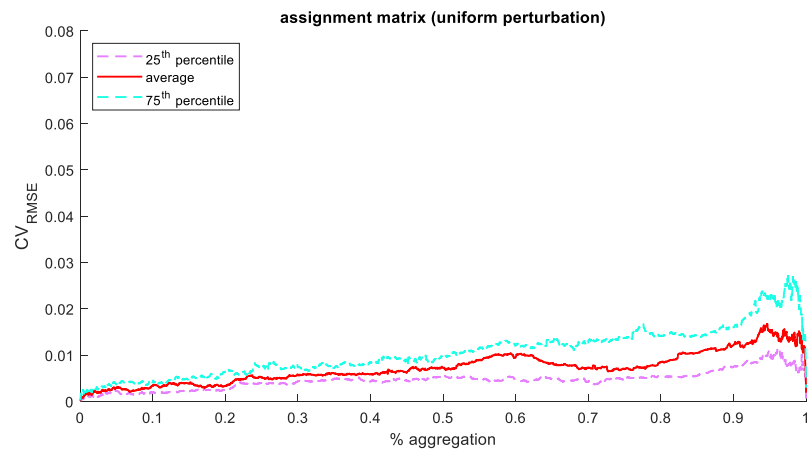
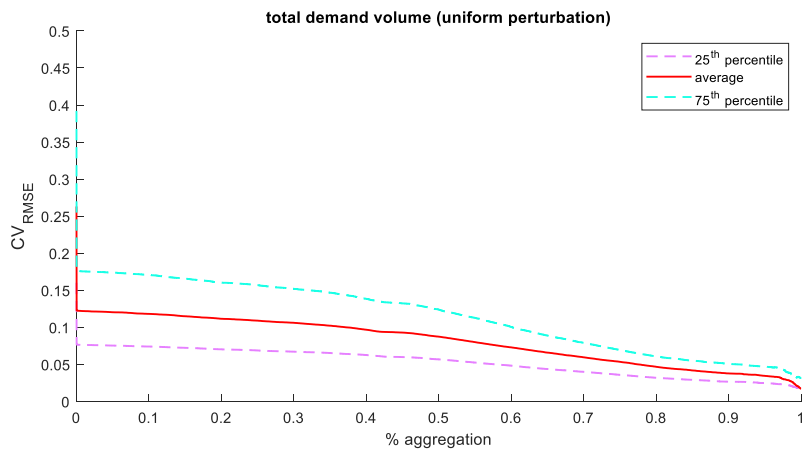
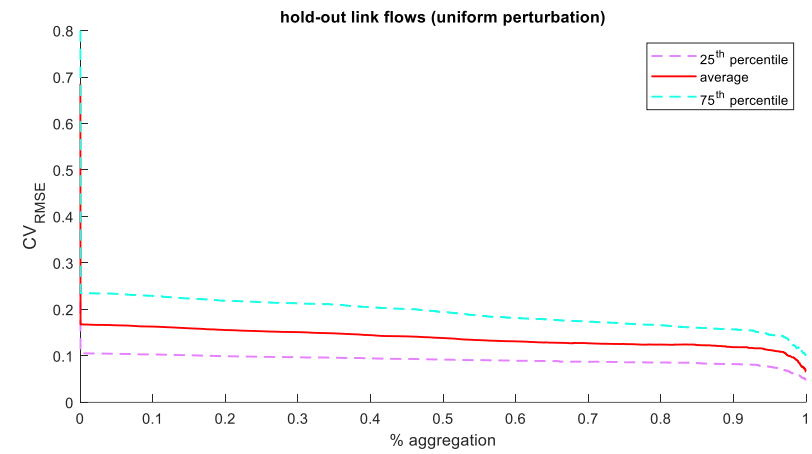
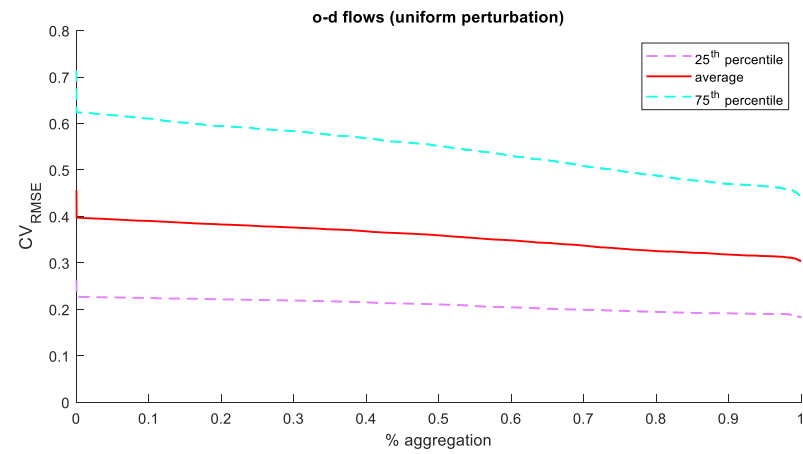


Fig. 3.10 Rete di Caserta: risultati sperimentali con perturbazione della stima a priori di tipo uniforme e 43 sezioni di conteggio

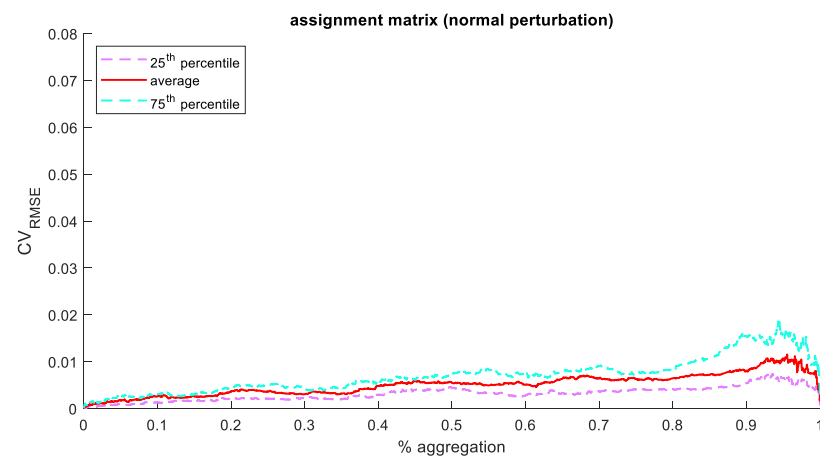
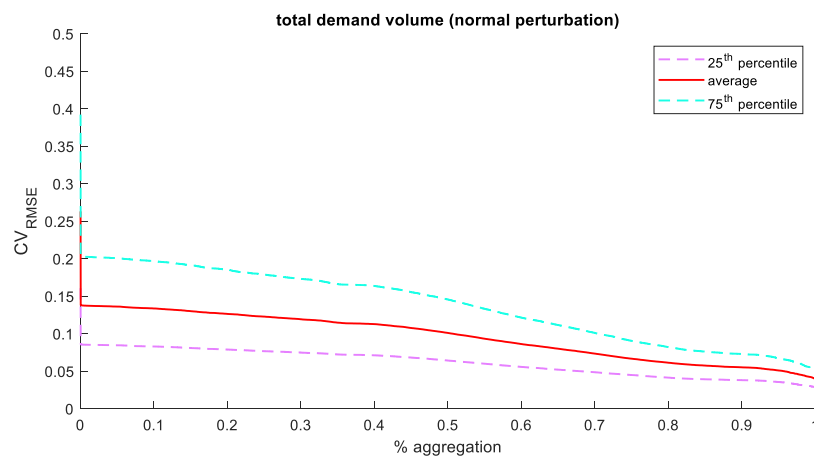
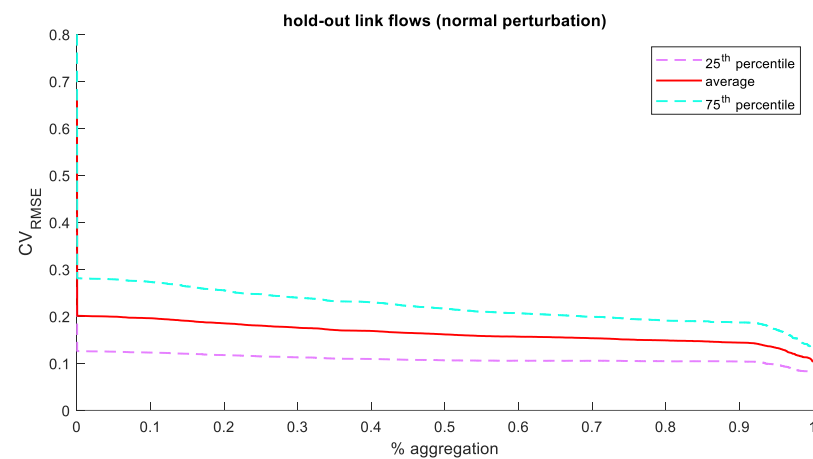
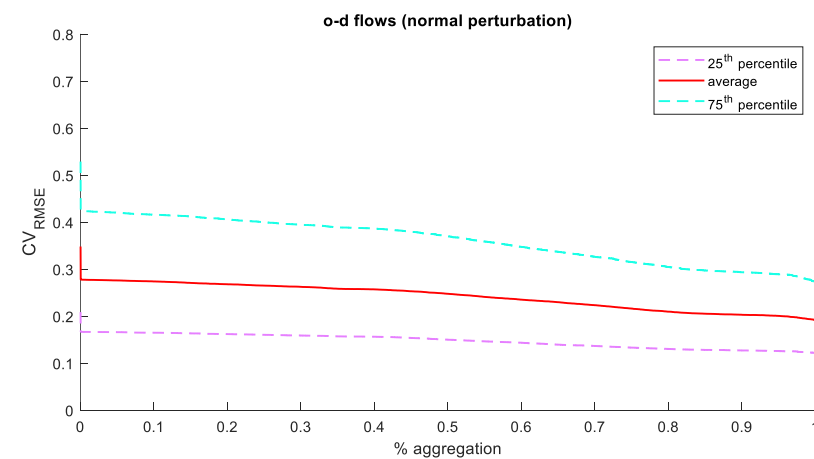


Fig. 3.11 Rete di Caserta: risultati sperimentali con perturbazione della stima a priori di tipo normale e 25 sezioni di conteggio

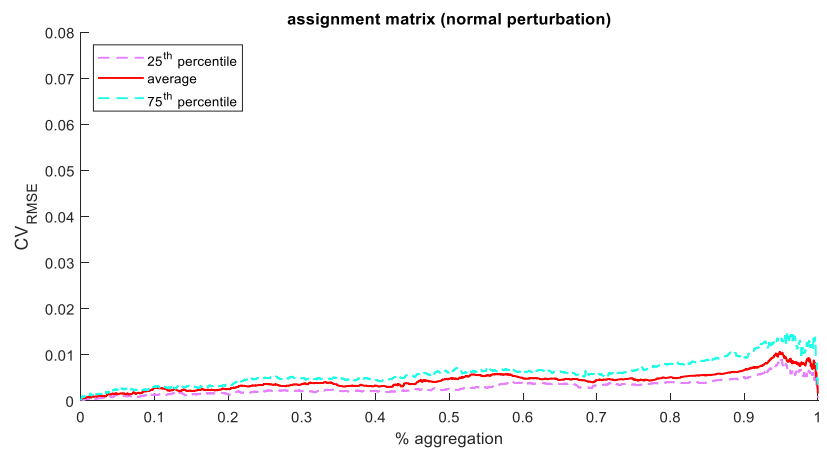
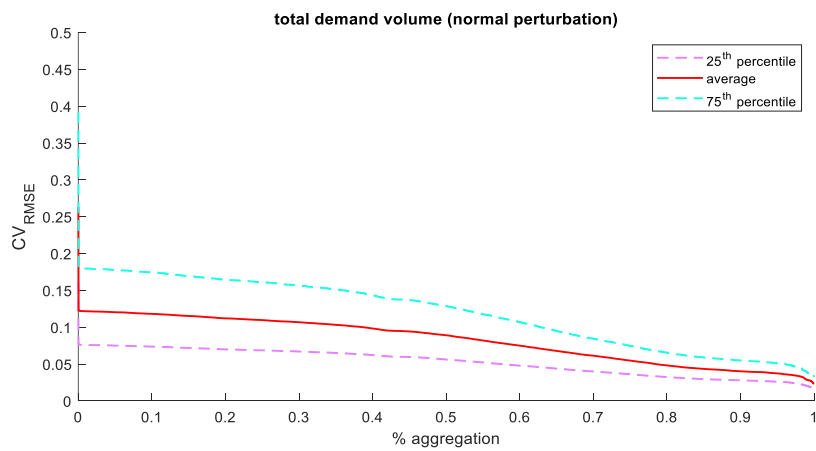
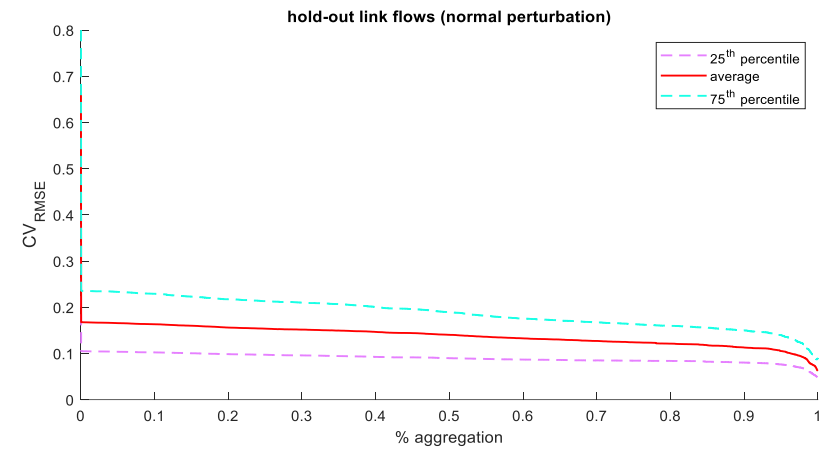
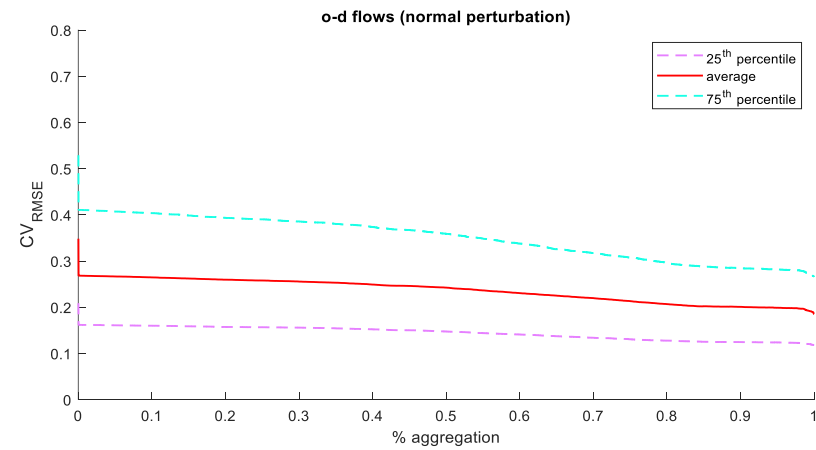


Fig. 3.12 Rete di Caserta: risultati sperimentali con perturbazione della stima a priori di tipo normale e 43 sezioni di conteggio

4. Correzione della matrice di domanda mediante clustering simultaneo di flussi *od*

Come anticipato alla fine del capitolo precedente, la metodologia descritta nel capitolo precedente nonostante si sia dimostrata efficace da un punto di vista di indicatori prestazionali, non risulta applicabile a reti di grandi dimensioni caratterizzate da un elevato numero di coppie *od*, in quanto richiederebbe la realizzazione di migliaia di procedure di correzione e quindi tempi di calcolo eccessivi. D'altra parte, l'applicazione di aggregazioni successive, in modo sequenziale, effettuata nel capitolo precedente, ha avuto soprattutto un obiettivo conoscitivo, e cioè di comprendere se e in che misura l'aggregazione di flussi *od* potesse migliorare l'efficacia della procedura di correzione di tali flussi.

Sulla base di tale analisi conoscitiva, in questo capitolo si propone una metodologia di correzione sempre basata sull'aggregazione, ma con delle tecniche di clustering che mirino a raggiungere direttamente un numero di coppie *od* risultanti in grado di bilanciare il numero delle equazioni disponibili, condizione che, come già ricordato precedentemente, Simonelli et al. 2012 hanno dimostrato essere necessaria per una efficace correzione della matrice *od*.

4.1. Descrizione metodologia

Come accennato alla fine della paragrafo precedente, la metodologia proposta in questo capitolo mira a clusterizzare le coppie *od* definite in fase di zonizzazione in modo da

raggiungere la condizione di osservabilità (Castillo et al., 2008). Essa si realizza quando il numero di equazioni linearmente indipendenti, rappresentate dalle sezioni di conteggio, è in numero maggiore o uguale alle incognite coppie od .

La procedura proposta richiede come input iniziali dell'algoritmo: una stima a priori vettore di domanda d^0 con cardinalità n_{od} coerente con un set iniziale di coppie od P^0 , un vettore di archi conteggiati C , una matrice di assegnazione M^{C0} , un numero finale di coppie od n_{od}^* a cui arrivare per raggiungere la condizione di osservabilità.

E' importante sottolineare che, date le dimensioni del problema, non è stato possibile utilizzare le classiche funzioni di clustering ma, è stato necessario proporre un algoritmo di clusterizzazione ad hoc, sufficientemente semplice da poter lavorare con un numero di incognite così elevato. Considerato che bisogna costruire un numero di cluster pari al numero degli archi conteggiati, la logica di tale algoritmo è quella di associare ad ogni arco conteggiato le coppie od che utilizzano maggiormente quell'arco, cioè con un maggior valore di m . Inoltre, nella procedura di clusterizzazione bisogna anche risolvere il problema dell'allocazione nei cluster di tutte le $n_{covnulla}$ coppie od con coverage nulla, cioè non intercettate da nessuna sezione di conteggio.

In sintesi, la struttura dell'algoritmo che implementa la procedura proposta è la seguente:

Step 1: Aggregazione finalizzata all'osservabilità

- a) Creazione di un numero di cluster pari al numero di sezioni di conteggio n_c ;
- b) Assegnazione ad ogni cluster di un numero di coppie od pari a $n_{elem}=n_{od}/n_c$

Per ogni arco conteggiato $c \in C$

- attribuzione a ciascun cluster di un numero di coppie od a coverage nulla pari a $n_{covnulla}/n_c$;
- Ordinamento decrescente delle coppie od in base al valore assunto dagli elementi della riga c della matrice d'assegnazione;
- attribuzione a ciascuno degli n_c cluster rappresentante un determinato arco conteggiato c , di un numero di coppie od pari a $n_{elem}-n_{covnulla}/n_c$ secondo l'ordinamento descritto al punto precedente;
- aggregazione del vettore di domanda e della matrice di assegnazione congruentemente con gli elementi contenuti nei cluster;

Step 2: Correzione della domanda

Correzione del vettore dei flussi od aggregati nello step 1 utilizzando i conteggi $c \in \mathbf{C}$, la matrice di assegnazione aggregata e lo stimatore GLS in modo da ottenere un vettore aggregato e corretto $\mathbf{d}^{GLS,*}$.

Step 3: Disaggregazione del flusso od corretto

Disaggregazione del vettore di domanda corretto $\mathbf{d}^{GLS}$ fino alla granularità iniziale costituita da P^0 coppie od .

L'aggregazione della domanda con riferimento al generico cluster k avviene utilizzando la seguente relazione:

$$d_k = \sum_{od \in k} d_{od}^{prior} \quad \forall k \quad (4.1)$$

L'aggregazione della matrice di assegnazione per il generico arco l e cluster k si definisce come:

$$m_k^l = \frac{\sum_{od \in k} m_{od}^l \cdot d_{od}^{prior}}{\sum_{od \in k} d_{od}^{prior}} \quad \forall k \quad (4.2)$$

Prima di eseguire la correzione, analogamente a quanto fatto nel capitolo 3, è stato necessario aggregare anche la matrice di dispersione.

La correzione della domanda aggregata è stata eseguita attraverso l'utilizzo del solito stimatore dei minimi quadrati, in particolare, mediante l'equazione di Aitken (2.25), la cui applicazione fornisce un nuovo vettore di domanda $\mathbf{d}^{GLS,*}$ avente k elementi $d_k^{GLS,*}$.

Dopo la procedura di correzione, la domanda aggregata-corretta deve essere riportata alla granularità iniziale attraverso una procedura di disaggregazione che utilizza come pesi la stima a priori del vettore di domanda. Con riferimento all' i -esima coppia od appartenente al set iniziale di coppie od P^0 :

$$d_{od_i}^{GLS} = d_k^{GLS,*} \cdot \frac{d_{od_i}^{prior}}{\sum_{od \in k} d_{od}^{prior}} \quad \forall od_i \in k, \forall k \quad (4.3)$$

Di seguito è riportato lo schema della procedura adottata:

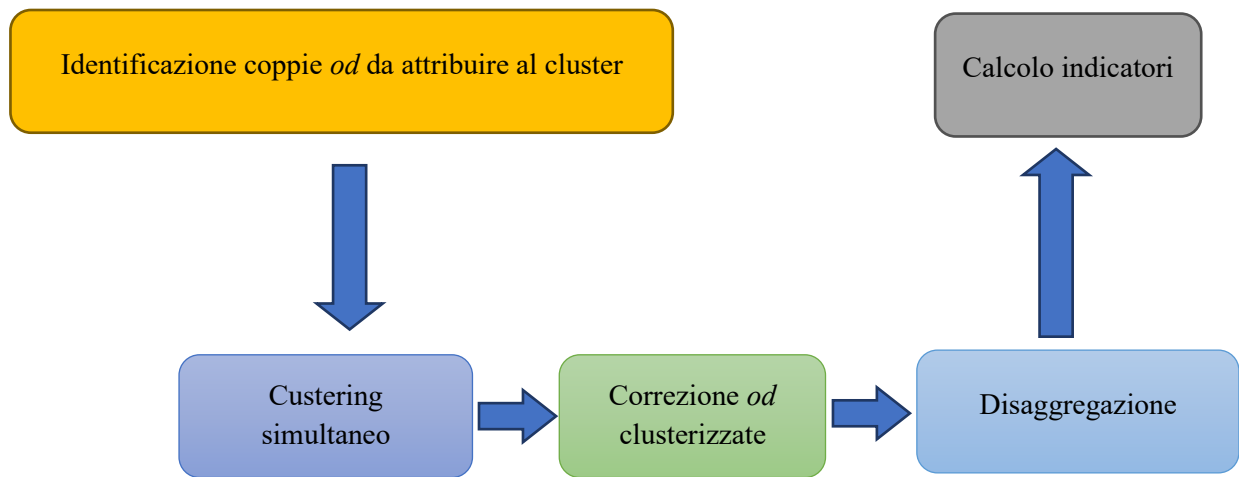


Fig. 4.1 Schema della metodologia proposta

Infine, a valle di ogni disaggregazione sono stati calcolati una serie di indicatori prestazionali, in particolare, lo stimatore utilizzato per la valutazione degli indicatori prestazionali è il $cvRMSE$ (3.4).

Anche in questo caso, le prestazioni della procedura di correzione proposta sono comparate con quelli ottenuti applicando la procedura di correzione senza effettuare nessuna clusterizzazione. In definitiva, quindi, si sono valutati gli indicatori prestazionali nei seguenti tre punti:

- (-1) nella condizione iniziale;
- (0) a valle di aver applicato la procedura standard di correzione GLS;
- (1) a valle di aver applicato la procedura di correzione proposta (clusterizzazione, correzione GLS, disaggregazione).

I punti -1 e 0 servono, al solito, a scopo comparativo, cioè a caratterizzare l'errore iniziale e l'errore ottenuto mediante l'applicazione della procedura di correzione classica della domanda.

4.2. Descrizione caso studio



Fig. 4.2 Rete di Torino,

L'esperimento di laboratorio è stato condotto sulla rete della città di Torino, costituita da 302 centroidi, circa, 20'000 archi e un numero di sezioni di conteggio pari a 300.

La metodologia descritta è stata applicata con due tipologie di setting degli esperimenti.

4.2.1. Setting 1

Il vettore di domanda ipotizzato vero $\mathbf{d}^{True}$ è stato ottenuto a partire dalla matrice di mobilità sistematica particella-particella fornita dell'ISTAT.

La stima a priori $\mathbf{d}^{Prior}$, invece, è stata ottenuta attraverso l'applicazione di un modello di domanda di tipo sequenziale, inoltre, per una maggiore generalità dei risultati, il modello è

stato applicato 36 volte cambiando i parametri di input, generando così 36 stime a priori del vettore di domanda.

Il modello sequenziale è stato specificato nel seguente modo in riferimento alla generica origine o e destinazione d

$$\begin{aligned}
 d_{od} &= d_o \cdot p[d|o] \\
 d_o &= \beta_{Emissione} \cdot Residenti_o \\
 V[d|o] &= \beta_{Attrazione} \cdot Addetti_d + \beta_{tempo} \cdot Tempo_{od} \\
 p[d|o] &= \frac{\exp(V[d|o])}{\sum_{d'} V[d'|o]}
 \end{aligned} \tag{4.4}$$

dove:

d_o rappresenta la domanda emessa dalla zona o , funzione dei residenti ($Residenti_o$) e di un coefficiente di emissione ($\beta_{Emissione}$)

$V[d|o]$ rappresenta l'utilità sistematica associata allo spostamento da o verso d ;

$p[d|o]$ rappresenta la probabilità di scelta della zona d come destinazione da parte dei flussi in uscita dalla zona o . Essa è stata stimata mediante l'applicazione di un modello di tipo logit multinomiale.

Il modello specificato (4.4) è in una forma molto semplificata, infatti, non è stata considerata la *logsum* sui modelli di scelta precedenti, ovvero, il modello di scelta modale.

Il valore assunto da residenti e addetti per ogni zona della rete è stato ricavato da database forniti dall'ISTAT, mentre l'insieme delle 36 coppie di coefficienti β , è stato ricavato combinando nella 4.4 6 possibili valori del coefficiente di emissione con 6 possibili valori del coefficiente di attrazione, partendo da valori ottenuti in studi precedenti:

$$\begin{aligned}
 \beta_{Emissione} &\in [2 \cdot 10^{-4} - 3 \cdot 10^{-4}], step 2 \cdot 10^{-4}; [persone^{-1}] \\
 \beta_{Attrazione} &\in [0.1 - 1.1], step 0.2; [persone^{-1}] \\
 \beta_{tempo} &-0.06 [minuti^{-1}]
 \end{aligned} \tag{4.5}$$

L'impedenza di trasporto tra le zone ($Tempo_{od}$) è stata stimata dal modello di offerta.

Per quanto riguarda la matrice di assegnazione, sono state formulate due diverse ipotesi che hanno consentito l'esplorazione di due tipologie di risultati:

- matrice d'assegnazione error-free, quindi matrice di assegnazione utilizzata nella procedura GLS coincidente con quella ipotizzata vera. In questo caso la matrice di assegnazione è stata ottenuta da un modello di scelta del percorso di tipo logit;
- matrice d'assegnazione non error-free, in questo caso si è ipotizzato una matrice di assegnazione vera ottenuta da un modello di scelta del percorso di tipo probit; mentre per la sua stima si è ipotizzata una matrice ottenuta da un modello di scelta del percorso di tipo logit, introducendo quindi così un errore di assegnazione.

4.2.2. Setting 2

Il vettore di domanda ipotizzato vero $\mathbf{d}^{True}$ è stato ottenuto sempre a partire dalla matrice di mobilità sistematica particella-particella fornita dell'ISTAT. Ad ogni elemento di tale vettore, corrispondente ad uno specifico individuo, è stata associata una traiettoria. Per far ciò, si è associata ad ogni individuo i una funzione di utilità di percorso espressa come una combinazione lineare di tempo di viaggio e costo monetario. Formalmente l'utilità percepita da ciascun individuo può essere definita nel seguente modo:

$$U_r^i = \beta_t^i \cdot t_r + \beta_c^i \cdot c_r + \varepsilon_r^i \quad (4.6)$$

Dove:

U_r^i è l'utilità del percorso r percepita dall'individuo i ;

t_r è il tempo viaggio sul percorso r ;

c_r è il costo monetario del percorso r ;

β_t^i, β_c^i sono i coefficienti della funzione di utilità relativa all'utente i ;

ε_r^i è il residuo aleatorio nell'utilità percepita dell'utente i sul percorso r .

I tempi di viaggio e i costi monetari sono stati desunti da un modello di offerta già a disposizione del gruppo di lavoro. I coefficienti del tempo di viaggio e del costo monetario sono estratti da due diverse distribuzioni normali monovariate. I valori medi di tali distribuzioni sono stati fissati in modo da avere un rapporto che sia uguale a 10 €/h (VOT) e le dispersioni delle distribuzioni di β_t e β_c sono state fissate in modo tale da non avere valori positivi irrealistici di essi all'interno della popolazione sintetica.

I residui aleatori ε_r^l sono stati calcolati in accordo con l'ipotesi di Daganzo e Sheffi (Daganzo & Sheffi, 1977) e cioè sommando residui aleatori di arco ε_l^l estratti da normali monovariate con la seguente distribuzione:

$$\varepsilon_l^l \sim N(0, \xi \cdot c_l) \quad (4.7)$$

dove c_l è il costo medio relativo all'arco l e ξ è un coefficiente di proporzionalità, valutato per ogni coppia od nel seguente modo:

$$\xi = cv^2 \cdot C_{od,min} \quad (4.8)$$

dove cv è il coefficiente di variazione fissato a seguito di studi parametrici, $C_{od,min}$ rappresenta il percorso di minimo costo per la generica coppia od .

Una volta generate tutte le traiettorie è possibile identificare la matrice di assegnazione vera e quella stimata per ogni valore di tasso di campionamento, calcolando il generico elemento $m_{\alpha,od}^l$ nel seguente modo.

$$m_{\alpha,od}^l = \frac{\Gamma_{\alpha,od}^l}{\Gamma_{\alpha,od}} \quad (4.9)$$

dove:

$\Gamma_{\alpha,od}^l$ rappresenta il numero di traiettorie, estratte con tasso di campionamento α , che connettono la coppia od passando per l'arco l .

$\Gamma_{\alpha,od}$ rappresenta il numero di traiettorie, estratte con tasso di campionamento α , che connettono la coppia od .

Per calcolare la matrice di assegnazione vera si deve utilizzare ovviamente l'equazione (4.9) con $\alpha=100\%$.

La stima a priori, $\mathbf{d}^{Prior}$, è stata ottenuta a partire da $\mathbf{d}^{True}$ con il metodo della stima diretta, ipotizzando diversi tassi di campionamento α , variabili da 1 a 100%. In ogni stima si è anche introdotto un errore di sovrastima della domanda del 20%, ipotizzando una non perfetta conoscenza dell'universo.

In corrispondenza ad ogni vettore di domanda campionato, e alle corrispondenti traiettorie, si è costruita anche una stima della mappa di assegnazione utilizzando sempre le 4.9, applicate al sottoinsieme delle traiettorie estratte. Per entrambi i setting lo stimatore utilizzato per la correzione della domanda è stato sempre il *GLS*, applicato mediante la sua formulazione in forma chiusa (2.6).

Analogamente a quanto ipotizzato nel capitolo 4, la matrice di dispersione della domanda Σ_d è stata assunta omoschedastica, diagonale e di due ordini di grandezza superiore alla matrice di dispersione dei flussi d'arco Σ_f , a sua volta ipotizzata diagonale, eteroschedastica con varianza proporzionale ai flussi d'arco conteggiati.

$\hat{f}_c$ è il vettore dei flussi d'arco conteggiati, ottenuto assegnando la domanda ipotizzata vera con la matrice di assegnazione ipotizzata vera.

4.3. Risultati e conclusioni

In questo paragrafo sono riportati i risultati degli esperimenti eseguiti sulla rete di Torino relativi alla metodologia proposta di correzione della domanda di trasporto con conteggi di traffico.

I grafici seguenti riportano il valore dell'indicatore *cvRMSE* relativo all'hold-out dei flussi di arco valutato nei seguenti tre punti:

- (-1) assegnando la stima a priori della domanda di trasporto;
- (0) assegnando la stima a priori corretta con la procedura di correzione GLS della domanda;
- (1) assegnando la domanda aggregata-corretta-disaggregata; l'aggregazione è eseguita fino al raggiungimento della condizione di osservabilità.

Nella figura (4.3) è rappresentato l'andamento dell'indicatore prestazionale *cvRMSE* sotto l'ipotesi di errore nullo sulla matrice di assegnazione, valutato rispetto all'hold-out dei flussi d'arco al variare della stima a priori della domanda da modello (setting 1 degli esperimenti). Come si evince dal grafico, effettuare la correzione della domanda a valle di averla aggregata in modo da raggiungere le condizioni di osservabilità (punto 1), introduce sempre un vantaggio molto significativo in termini di efficacia della procedura e di qualità dell'output, indipendentemente dalla stima a priori della domanda.

Le 36 stima a priori della domanda assegnate alla rete, generano dei flussi d'arco a cui è associato un indicatore d'errore variabile tra 0.85 e 2.50 nella condizione iniziale (punto-1). A seguito dell'applicazione della procedura di correzione classica della domanda, lo stesso errore ha una variabilità più stretta e precisamente tra 0.70 e 2.25 (punto 0). Infine, grazie alla procedura proposta si ottiene un errore con variabilità compresa tra 0.55 e 0.99. La

procedura proposta contenente dei miglioramenti percentuali compresi tra il 56% (da 2.25 a 0.99) e il 21 % (da 0.7 a 0.55).

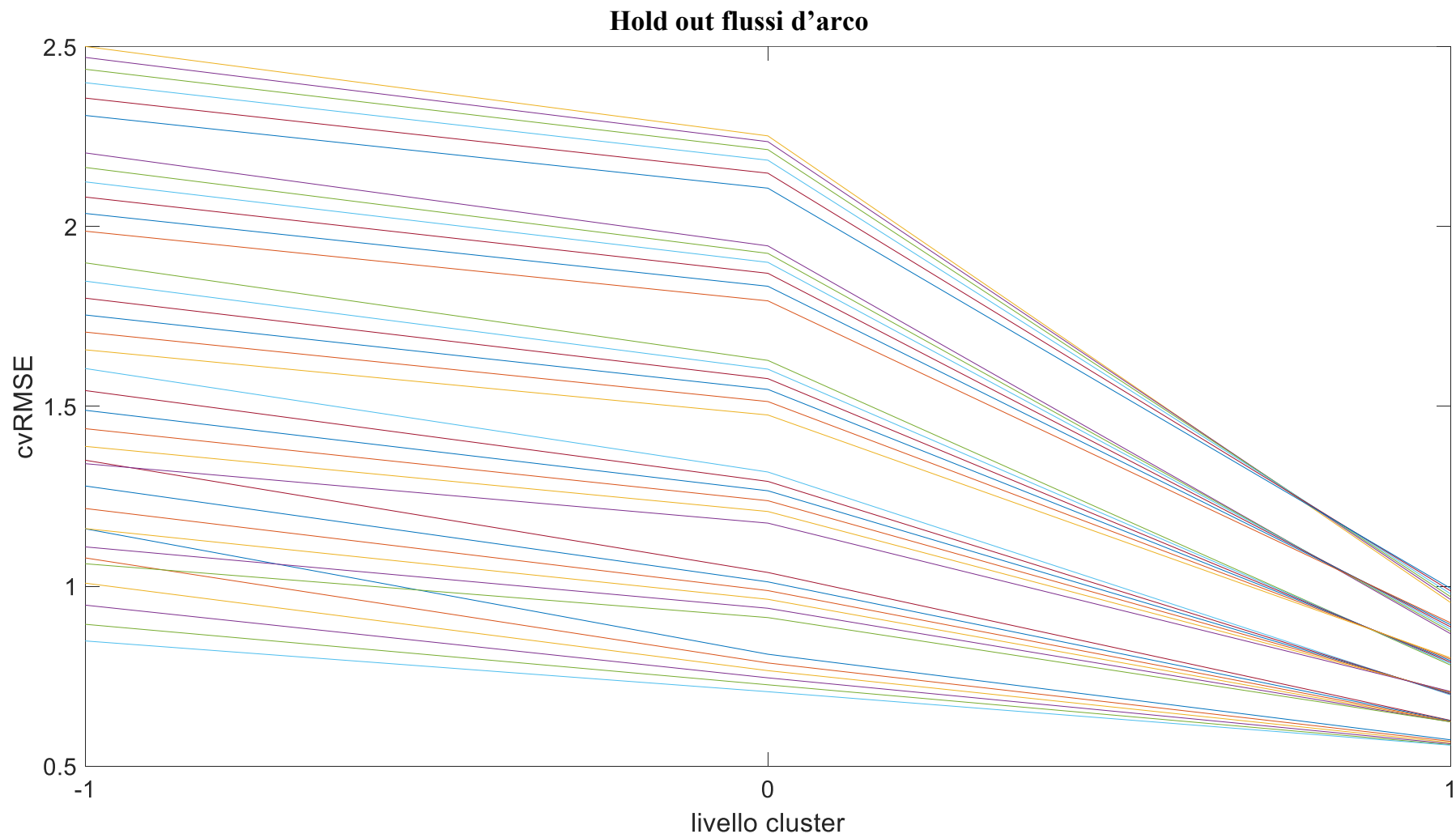


Fig. 4.3 cvRMSE hold out flussi d'arco (setting 1) al variare delle stime a priori della domanda, senza introduzione di un errore in fase di assegnazione.

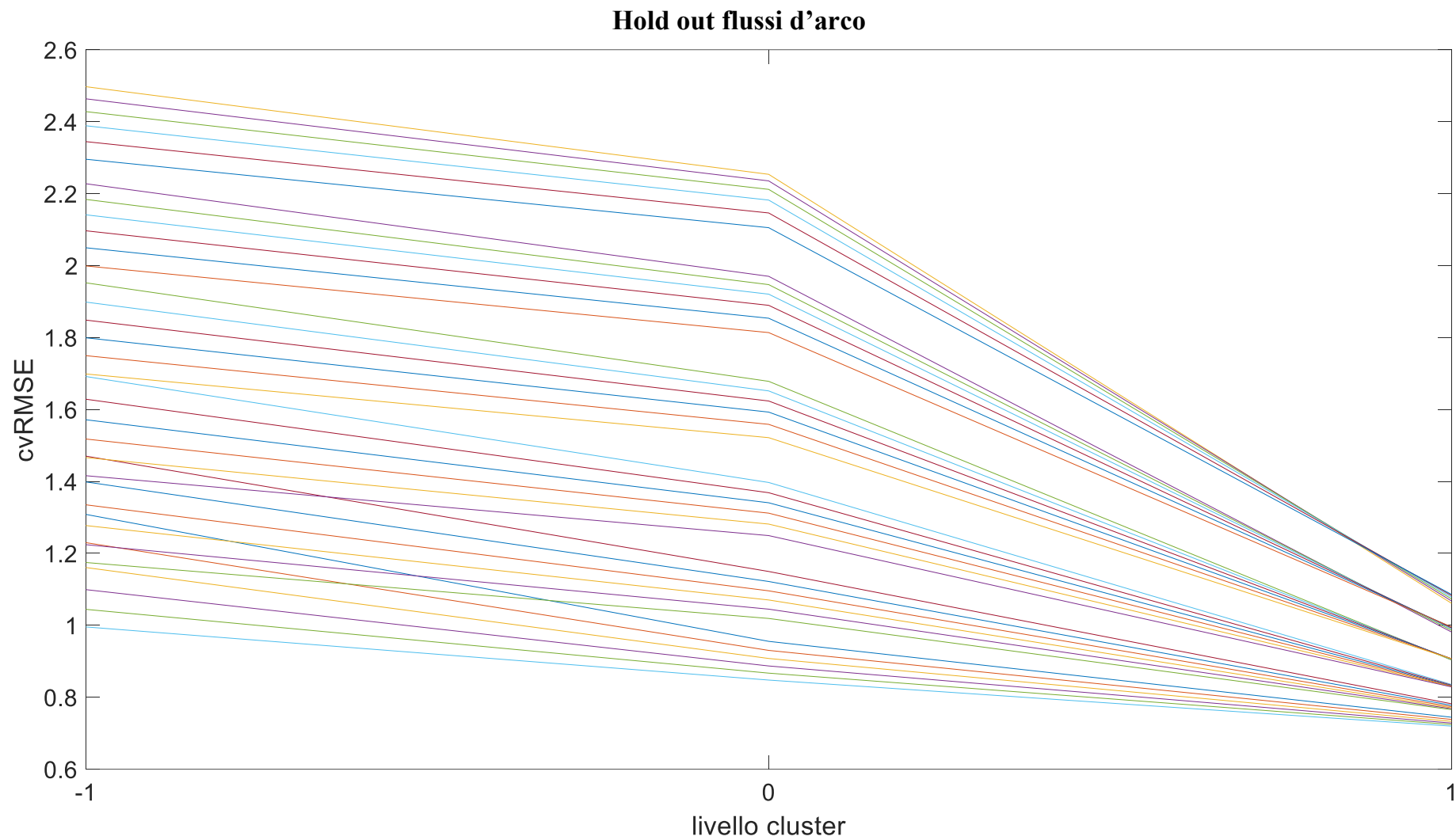


Fig. 4.4 cvRMSE hold out flussi d'arco (setting 1) al variare delle stime a priori della domanda, con introduzione di un errore in fase di assegnazione.

Nella figura (Fig. 4.4) sono riportati i risultati degli esperimenti condotti con il setting 1 in cui è stato introdotto un errore nella stima della matrice di assegnazione. Aver introdotto un errore nella matrice di assegnazione comporta, ovviamente, una peggiore qualità dell'output rispetto al caso con assenza di errore di assegnazione (Fig. 4.3). Nonostante l'errore, però, le curve hanno tutte un andamento decrescente, con l'ottimo raggiunto in corrispondenza dell'applicazione della procedura proposta.

Le 36 stima a priori della domanda assegnate alla rete, generano dei flussi d'arco a cui è associato un indicatore d'errore variabile tra 0.99 e 2.49 nella condizione iniziale (punto-1). A seguito dell'applicazione della procedura di correzione classica della domanda lo stesso errore ha una variabilità più stretta e precisamente tra 0.85 e 2.25 (punto 0). Infine, grazie alla procedura proposta si ottiene un errore compreso tra 0.71 e 1.08 (punto 1). Dunque, la procedura proposta introduce dei miglioramenti percentuali compresi tra il 52% (da 2.25 a 1.08) e il 16% (da 0.85 a 0.71).

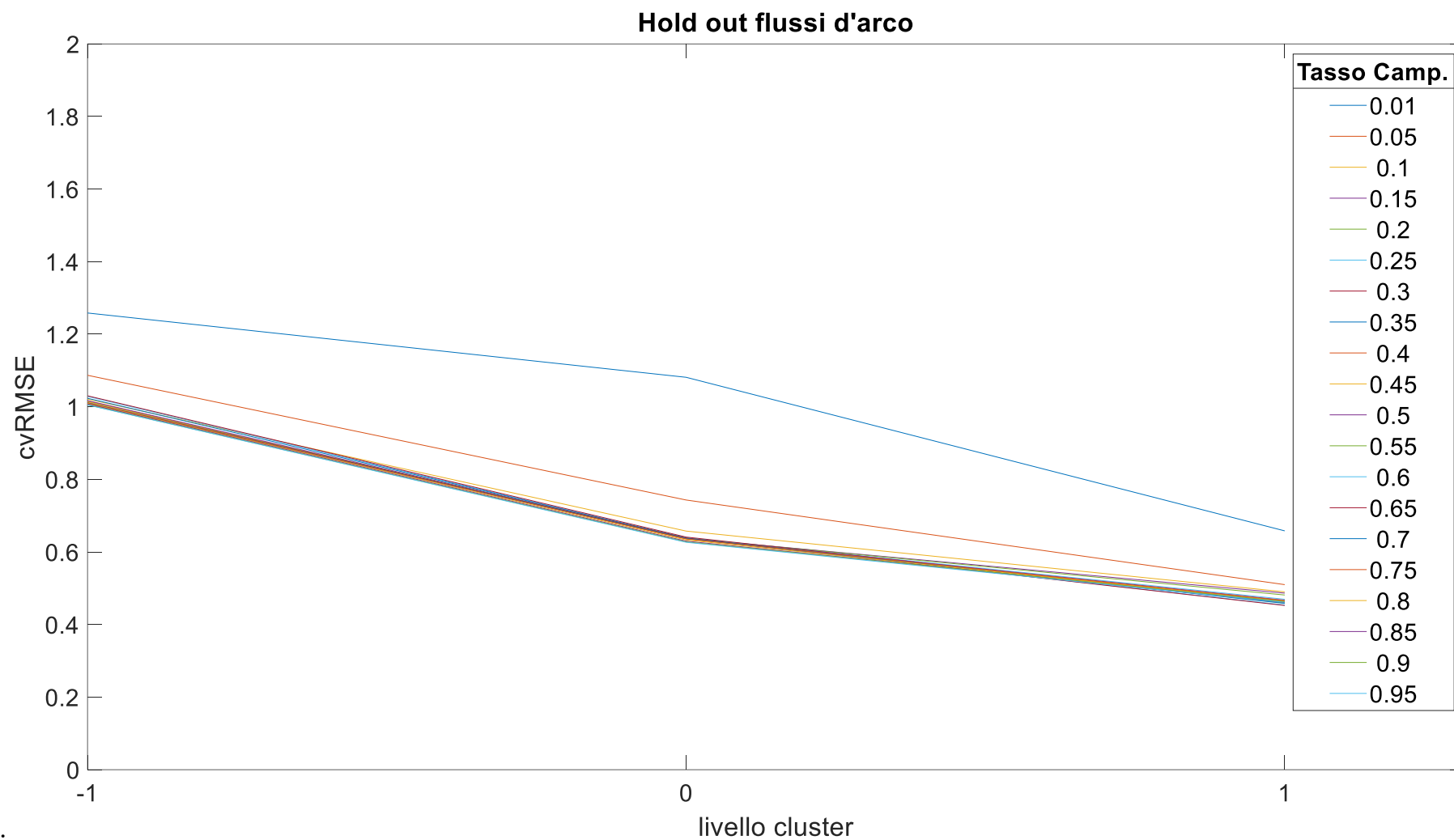


Fig. 4.5 cvRMSE hold out flussi d'arco (Setting 2) al variare del tasso di campionamento

Nella figura (Fig. 4.5) sono riportati i risultati della sperimentazione con riferimento al setting 2. In particolare, è riportato l'andamento dell'indicatore d'errore $cvRMSE$ valutato rispetto all'hold out dei flussi d'arco per diversi tassi di campionamento. Anche in questo caso, si può osservare un miglioramento di tale indicatore, effettuando la correzione nella condizione di osservabilità, con una riduzione mediamente pari al 28% (da 0.63 nel punto 0 a 0.45 nel punto 1). Si osservi che la correzione non è mai perfetta, in virtù dell'ipotesi fatta sulla non corretta conoscenza del valore della domanda complessiva.

5. Correzione della matrice di domanda mediante modello gravitazionale q-generalizzato e autocorrelazioni spaziali basate sulle adiacenze

Nel presente capitolo di tesi è illustrata un'ulteriore metodologia di correzione della domanda di trasporto con conteggi di traffico. In particolare, è stata investigata la possibilità di legare la stima della domanda di trasporto ad un modello di tipo q-gravitazionale generalizzato sfruttando la q-algebra e le sue proprietà (Box & Cox, 1964; Tsallis, 1988; Umarov et al., 2008), esempi applicati al campo dei trasporti sono riportati in diversi articoli scientifici (Tinessa, 2021; Tinessa et al., 2020). Essa risulta essere una particolare algebra deformata che consente di generalizzare alcune delle procedure già presenti in letteratura come i modelli gravitazionali di tipo log-lineare semplici e procedure analoghe all'analisi fattoriale.

La logica del modello proposto, analogamente ad un modello gravitazionale, è che un flusso *od* può essere spiegato attraverso delle variabili legate alla zona di origine, alla zona di destinazione e all'impedenza di trasporto che separa le due zone. Mentre in un classico modello gravitazionale le variabili legate a ciascuna zona sono note e i coefficienti del modello sono da stimare, al contrario, nell'approccio qui proposto le variabili legate alle

zone non sono delle quantità note, bensì anch'essa da stimare insieme ai coefficienti del modello.

Il modello proposto si basa su un'assunzione, cioè, se un flusso emesso da una certa zona di origine assume un valore elevato, probabilmente anche le zone limitrofe avranno un flusso emesso alto, con analoga situazione nella zona di destinazione. Quanto appena esposto si traduce nell'ammettere che esistano delle correlazioni spaziali tra le zone o in generale che la variabile flusso od è spazialmente autocorrelata. La presenza di correlazione è funzione della tipologia della rete e delle masse presenti nelle zone (residenti, addetti, industrie, etc.), in caso di zone a destinazione d'uso simili ci si aspetta un segno positivo del parametro di correlazione, viceversa la correlazione non è attesa essere significativa.

Chiaramente, fissare a priori l'esistenza e l'entità delle auto-correlazioni potrebbe essere una forzatura e quindi rivelarsi un'ipotesi sbagliata. Per questo motivo sono i dati a guidare la ricerca delle autocorrelazioni attraverso un processo di ottimizzazione. Infine, dopo aver stimato il parametro di correlazione è necessario procedere con un test statistico per valutarne la significatività. Secondo l'approccio proposto, il grado ed il numero di correlazioni attive sono definiti attraverso i coefficienti β_{o_i} β_{d_i} del modello, esprimibili attraverso la seguente espressione:

$$\begin{aligned}
 \ln_q(d_{od}) &= \sum_{i=1}^N \beta_{o_i} \cdot \ln_q(x_{o_i}) + \sum_{i=1}^N \beta_{d_i} \cdot \ln_q(x_{d_i}) \\
 &\quad + \sum_{j=1}^{\gamma} \beta_m \cdot \ln_q(x_m) + \beta_c \cdot \ln_q(\Gamma_{od}) + \ln_q(\xi_{od}) \\
 &= \sum_{i=1}^N \beta_{o_i} \cdot \ln_q(x_{o_i}) + \sum_{i=1}^N \beta_{d_i} \cdot \ln_q(x_{d_i}) \\
 &\quad + \sum_{m=1}^{\gamma} \beta_m \cdot \ln_q(x_{m_j}) + \beta_c \cdot \ln_q(\Gamma_{od}) + \xi_{od}^*
 \end{aligned} \tag{5.1}$$

d_{od} rappresenta la domanda di mobilità avente origine o e destinazione d ;

o e d rappresentano rispettivamente le zone di origine e destinazione degli spostamenti;

Γ_{od} rappresenta un'impedenza di trasporto posta pari a $\exp(t_{od})$;

t_{od} rappresenta distanza misurata come tempo tra tutte le coppie od ;

ξ_{od}^* rappresenta il termine di errore introdotto dal modello;

N rappresenta il numero di zone in cui è stata discretizzata l'area di studio.

x_{o_i}, x_{d_i} sono le incognite della procedura legate alla i -esima origine o_i e destinazione d_i ;

x_m sono delle ulteriori incognite del modello la cui numerosità è pari a γ ;

$\beta_{o_i}, \beta_{d_i}, \beta_m$ sono i coefficienti del modello;

β_c è un coefficiente del modello legato all'impedenza di trasporto da stimare, dunque è incognito. Di conseguenza il numero totale di incognite è pari a $2N + \gamma + 1$.

Analogamente alle regole dell'algebra classica, il q-prodotto può essere trasformato in una q-somma mediante trasformazione q-logaritmica:

La precedente equazione (5.1) è equivalente ad una regressione di box-cox con parametro $1-q$, in cui l'operatore q-logaritmo per la generica variabile ψ può essere descritto attraverso l'equazione (5.2). La q-algebra può essere utilizzata per capire quanto il modello di regressione descritto dall'equazione (5.1) si avvicina ad una tipologia lineare ($q = 0$) o moltiplicativa ($q \rightarrow 1$).

$$\ln_q(\psi) = \begin{cases} \frac{\psi^{1-q} - 1}{1-q}; & \text{per } q \neq 1 \\ \ln(\psi); & \text{per } q = 1 \end{cases} \quad (5.2)$$

L'equazione (5.3) può essere espressa in forma vettoriale nel seguente modo:

$$\mathbf{d}^* = \mathbf{B} \cdot \mathbf{x}^* + \boldsymbol{\xi} \quad (5.3)$$

dove:

$\mathbf{d}^*$ rappresenta il q-logaritmo del vettore di domanda;

$\mathbf{x}^*$ rappresenta il q-logaritmo delle variabili ricercate da cui dipende il modello;

$\boldsymbol{\xi}$ rappresenta il vettore degli errori introdotti dal modello;

$\mathbf{B}$ rappresenta la matrice dei coefficienti del modello avente dimensione $N^2 \times (2N + \gamma + 1)$.

5.1. Descrizione metodologia

L'obiettivo della metodologia descritta in questo capitolo è ottimizzare il vettore di variabili $\mathbf{x}$, il parametro q di un modello q-gravitazionale generalizzato e stabilire il valore dei coefficienti della matrice $\mathbf{B}$, in modo da generare una stima della domanda di trasporto, la

quale assegnata alla rete fornirà dei flussi d'arco compatibili con quelli osservati nelle sezioni di conteggio.

Per testare l'efficacia della metodologia proposta, analogamente a quanto fatto nel capitolo 4, sono stati valutati degli indicatori prestazionali in corrispondenza:

- (-1) della condizione iniziale
- (0) a valle della correzione effettuata con la procedura GLS classica
- (1) a valle della correzione effettuata applicando la metodologia proposta

La metodologia proposta prevede di stimare il valore dei parametri incogniti del modello $\mathbf{x}$ la cui numerosità è $2N + \gamma + 1$ e del parametro q attraverso un processo di ottimizzazione che sfrutta i conteggi di traffico ed una stima a priori della domanda di trasporto. Dunque, il numero complessivo di incognite da stimare è $2N + \gamma + 2$. Inoltre, è necessario identificare un criterio per la scelta dei coefficienti presenti nella matrice $\mathbf{B}$ (5.3).

È importante sottolineare che una parte⁵ dei coefficienti presenti all'interno della matrice $\mathbf{B}$ definiscono il grado e il numero di correlazioni che si attiva tra le zone dell'area di studio. Per questo motivo sono state proposte diverse tecniche per la costruzione della matrice $\mathbf{B}$ in modo da tener conto delle correlazioni spaziali.

Una prima metodologia di autocorrelazioni spaziali si basa sull'introduzione di un valore di soglia, definito *bandwidth* b (Fotheringham et al., 2002), oltre il quale non si considera più alcuna correlazione. Secondo tale approccio, il valore assunto dai coefficienti nella matrice $\mathbf{B}$ è esprimibile attraverso una funzione biquadratica, inversamente proporzionale al tempo di percorrenza tra le coppie *od*.

Nel caso in cui il numero di incognite costituenti il vettore $\mathbf{x}$ sia pari a $2N+2$, quindi $\gamma=0$, gli elementi della matrice $\mathbf{B}$ possono essere calcolati nel seguente modo:

⁵ Contribuiscono alle correlazioni solo i coefficienti β_{oi} , β_{di}

$$\begin{aligned}
\beta_{(o_i d_k), j'} &= \begin{cases} 1; \text{ se } i = j' \\ \left(1 - \left(\frac{t_{o_i j'}}{b}\right)^2\right)^2; \text{ se } i \neq j' \text{ and } t_{o_i j'} < b \\ 0 \text{ altrimenti} \end{cases} \\
\beta_{(o_i d_k), N+j''} &= \begin{cases} 1; \text{ se } k = j'' \\ \left(1 - \left(\frac{t_{d_k j''}}{b}\right)^2\right)^2; \text{ se } k \neq j'' \text{ and } t_{d_k j''} < b \\ 0 \text{ altrimenti} \end{cases} \\
\beta_{(o_i d_k), 2N+1} &= \exp(t_{o_i d_k}) \\
\text{con } j' \text{ e } j'' &\in [1, N]; \\
o_i, d_k &\in [1, N];
\end{aligned} \tag{5.4}$$

dove:

$t_{o_i j'}$ rappresenta il tempo di percorrenza tra la zona o_i e j' ;

$t_{d_k j''}$ rappresenta il tempo di percorrenza tra la zona d_k e j'' ;

$t_{o_i d_k}$ rappresenta il tempo di percorrenza tra la zona o_i e d_k

b rappresenta il valore di bandwidth che a sua volta può essere ottenuto attraverso una procedura di ottimizzazione ovvero fissato a priori.

Un approccio alternativo a quello descritto dalla relazione (5.4) per la definizione dei coefficienti della matrice $\mathbf{B}$ si basata sull'adiacenza topologica tra le zone Ω . Infatti, il criterio proposto di seguito (5.5) definisce un coefficiente di correlazione ρ da cui dipendono i termini della matrice $\mathbf{B}$.

$$\begin{aligned}
\beta_{(o_i d_k), j'} &= \begin{cases} 1; \text{ se } i = j' \\ \rho; \text{ se } i \neq j' \text{ and } o_{j'} \in \Omega_{o_i} \\ 0 \text{ altrimenti} \end{cases} \\
\beta_{(o_i d_k), N+j''} &= \begin{cases} 1; \text{ se } k = j'' \\ \rho; \text{ se } k \neq j'' \text{ and } d_{j''} \in \Omega_{d_k} \\ 0 \text{ altrimenti} \end{cases} \\
\beta_{(o_i d_k), 2N+1} &= \exp(t_{o_i d_k}) \\
o_i, d_k &\in [1, N]; \\
\text{con } j' \text{ e } j'' &\in [1, N];
\end{aligned} \tag{5.5}$$

Dove Ω_{o_i} e Ω_{d_k} rappresentano rispettivamente l'insieme delle zone adiacenti alla zona o_i e alla zona d_k . Anche in questo caso ρ può essere ottenuto attraverso una procedura di ottimizzazione oppure fissato a priori.

Nel caso in cui γ sia diverso da zero, la matrice $\mathbf{B}$ assumerà dimensione $N^2 \times (2N+1+\gamma)$. In questo caso la matrice $\mathbf{B}$ può essere scritta come il concatenamento in orizzontale della matrice $\mathbf{B}$ definita nel caso (5.4, 5.5) e di una matrice $\mathbf{B}_\gamma$ il cui generico elemento può essere calcolato nel seguente modo.

$$\beta_{\gamma(o_i d_k), g} = \begin{cases} 1; & \text{se } o_i d_k \in Cluster_g \\ 0 & \text{altrimenti} \end{cases}$$

$$g \in [1, \gamma]$$

$$o_i, d_k \in [1, N];$$
(5.6)

dove $Cluster_g$ rappresenta la g -esima partizione dell'insieme delle coppie od , ottenuta attraverso l'applicazione di un algoritmo di clustering sulla variabile tempo di viaggio tra le coppie od . Le ulteriori variabili γ servono quindi a tener conto della variabile esplicativa tempo di percorrenza od nella formazione di un flusso od .

Un'altra tecnica esplorata per la costruzione della matrice $\mathbf{B}$, si basa sul considerare i soli flussi od rilevanti. Essi rappresentano quelle particolari coppie od che danno il maggior contributo alla formazione dei flussi sugli archi principali⁶ della rete. Secondo tale criterio, sono stati esclusi (quindi posti a zero) i coefficienti della matrice $\mathbf{B}$ non associati alle coppie od rilevanti. Le coppie od rilevanti sono state identificate in base ad un ordinamento decrescente della stima a priori della domanda di trasporto, andando a selezionare le prime Φ coppie od , con Φ variabile parametricamente. E' evidente che al diminuire di Φ , da un lato diminuiscono le variabili da dover stimare - e quindi la procedura di correzione dovrebbe diventare più efficace - e dall'altro si trascurano un numero sempre crescente di flussi od che contribuiscono alla formazione dei flussi di arco, e quindi si introduce un errore nella stima di questi ultimi. L'obiettivo, quindi al solito, è quello di provare ad individuare il valore di Φ di miglior compromesso (trade off).

⁶L'importanza di un arco è definita in base alla gerarchia delle strade (autostrade, extraurbane principali, etc.).

Dunque, la matrice $\mathbf{B}$ può essere intesa come una funzione di diversi parametri b , ρ , γ e Φ .

$$\mathbf{B} = \mathbf{B}(\rho, b, \gamma, \Phi) \quad (5.7)$$

Una volta identificato il criterio di calcolo dei coefficienti della matrice $\mathbf{B}$, è necessario eseguire un'ottimizzazione delle incognite $\mathbf{x}, q$ e $\mathbf{B}(\rho, b, \gamma, \Phi)$ del modello, utilizzando come dati osservati i flussi d'arco nelle sezioni di conteggio stradali $\hat{\mathbf{f}}_c$ e la stima a priori della domanda $\hat{\mathbf{d}}$.

$$\begin{aligned} (\mathbf{x}, q, \rho, b, \gamma, \Phi)_{NLS} = \operatorname{argmin} \{ & (\mathbf{M}^c \cdot \mathbf{B}(\rho, b, \gamma, \Phi) \cdot \mathbf{x} - \hat{\mathbf{f}}_c)^T \cdot \boldsymbol{\Sigma}_f^{-1} \\ & \cdot (\mathbf{M}^c \cdot \mathbf{B}(\rho, b, \gamma, \Phi) \cdot \mathbf{x} - \hat{\mathbf{f}}_c) \\ & + (\mathbf{B}(\rho, b, \gamma, \Phi) \cdot \mathbf{x} - \hat{\mathbf{d}})^T \cdot \boldsymbol{\Sigma}_d^{-1} \\ & \cdot (\mathbf{B}(\rho, b, \gamma, \Phi) \cdot \mathbf{x} - \hat{\mathbf{d}}) \} \end{aligned} \quad (5.8)$$

dove:

$\mathbf{M}^c$ rappresenta la matrice di assegnazione relativa alle sole sezioni di conteggio;

$\boldsymbol{\Sigma}_f$ rappresenta la matrice di dispersione dei flussi d'arco;

$\boldsymbol{\Sigma}_d$ rappresenta la matrice di dispersione della domanda;

$\hat{\mathbf{d}}$ è la stima a priori della domanda di trasporto.

L'ultima fase della procedura consiste nel calcolo degli indicatori prestazionali valutati rispetto ai flussi d'arco, anche in questo caso lo stimatore utilizzato è il *cvRMSE*.

5.2. Descrizione caso studio

La procedura proposta in questo capitolo è stata nuovamente testata sulla rete di Torino (Fig.4.2) già descritta nel capitolo 4, costituita da 302 centroidi, circa 20'000 archi e un numero di sezioni di conteggio pari a 300.

Per testare l'efficacia della procedura proposta di correzione della domanda con conteggi di traffico, è stato necessario progettare una campagna sperimentale, in cui si è ipotizzato di far variare una serie di grandezze e di fissarne altre. Analogamente al capitolo 4, è stata fissata la domanda ipotizzata vera, proveniente dalla mobilità sistematica particella-particella ISTAT e i flussi d'arco veri sono stati ottenuti assegnando la precedente domanda con una matrice d'assegnazione ottenuta da modello di scelta del percorso di tipo logit a

doppio passo. Per quanto riguarda le ipotesi sulle stime a priori della domanda e le relative matrici di assegnazione si rimanda al capitolo 4.

La numerosità dei parametri che è possibile far variare nella procedura proposta (stima a priori della domanda, numero di coppie *od* variabili Φ , numero di sezioni di conteggio, q , γ , ρ , b) definisce un piano sperimentale molto vasto ed estremamente lungo da un punti di vista computazionale, visti anche i tempi di calcolo necessari per far girare la metodologia proposta, soprattutto in alcuni contesti sperimentali (attivazione di tutte le variabili). Di conseguenza, sono state esplorate solo alcune delle possibili combinazioni di parametri, cercando di ottenere dei risultati dai quali dedurre delle considerazioni di carattere il più possibile generale. La realizzazione del piano sperimentale completo è lasciata quindi a sviluppi futuri.

In particolare, sono state testate più stime a priori della domanda per ogni parametro, ma non tutte le possibili intersezioni di parametri e stime a priori.

Il parametro q è stato testato nell'intervallo 0-1; ρ e b inizialmente sono stati posti pari a zero, successivamente ottimizzati; per γ si sono ipotizzati tre valori: 0, 50 e 100; Φ è stato testato nell'intervallo 5mila-45mila.

In merito a Φ , è importante sottolineare che la matrice di domanda, relativa al caso studio di Torino, è costituita da 90mila elementi di cui la metà circa hanno un valore nullo del flusso. Di conseguenza, i risultati che si ottengono applicando la procedura con Φ pari a 45mila sono pressoché equivalenti a quelli ottenibili con valori superiori di tale parametro.

Nella seguente tabella (5.1) è riportato lo schema degli esperimenti condotti al variare dei parametri definiti precedentemente. In particolare, i numeri presente all'interno della tabella in carattere corsivo indicano quante stime a priori della domanda sono state utilizzate per eseguire quello specifico esperimento.

In particolare, quando tale numero è di colore nero, vuol dire che le stime a priori della domanda, sono state generate mediante un modello di domanda, variando il valore dei suoi coefficienti; il colore blu, invece, indica delle stime a priori della domanda, ottenute mediante procedura di stima diretta facendo variare il tasso di campionamento. In entrambi i casi, il numero riportato vicino, in corsivo tra parentesi tonde, indica la figura dove sono riportati i risultati di tale sperimentazione.

Parametri della procedura

Coppie <i>od</i> variabili Φ :												
tutte (90k)						rilevanti (5k-45k)						
Numero di conteggi:	300			Variabili (50-700)			300			Variabili (50-700)		
γ :	0	50	100	0	50	100	0	50	100	0	50	100
$q=\rho=b=0$:	36 (5.1)	6 (5.4)	6 (5.5)	1 (5.3)			1 (5.2)	1 (5.6)	1 (5.7)			
	1 (5.8)						10 *					
$q=\forall, \rho=b=0$:	1 (5.19)											

Tab. 5.1

numero di stime a priori della domanda utilizzate, variabili al variare dei coefficienti del modello

numero di stime a priori della domanda utilizzate, variabili al variare del tasso di campionamento

* Risultati riportati dalla figura (5.9) alla figura (5.18)

5.3. Risultati e conclusioni

Nel presente paragrafo sono riportati i risultati del piano sperimentale condotto, la cui descrizione sintetica è riportata nella tabella 5.1.

Il racconto degli esperimenti segue un filo conduttore associato al valore di q , infatti, saranno prima descritti tutti gli esperimenti per i quali q assume il valore 0, vale a dire modello di regressione (5.1) di tipo additivo. Successivamente saranno inseriti i risultati con q variabile parametricamente, infine i risultati con q ottenuto attraverso un processo di ottimizzazione.

In tutte le figure di seguito riportate è rappresentato l'andamento dell'indicatore d'errore $cvRMSE$ valutato in tre punti specifici rispetto all'hold out dei flussi d'arco quando $\gamma=0$ e in quattro specifici punti per $\gamma \neq 0$. L'indicatore è stato valutato confrontando, al solito, il vettore dei flussi d'arco veri con quello ottenuto assegnando:

- punto (-1): la stima a priori della domanda;
- punto (0): la domanda corretta mediante la procedura di correzione GLS classica;
- punto (1) la domanda corretta con la procedura proposta imponendo $\gamma=0$;
- punto (2); la domanda corretta con la procedura proposta e ottimizzando anche le variabili γ . Questo punto si attiva, ovviamente, solo nelle sperimentazioni in cui sono state introdotte anche le variabili γ . Tale punto (2), inserito all'interno dello stesso grafico, consente di evidenziare lo specifico miglioramento ottenibile con l'aggiunta delle sole variabili γ .

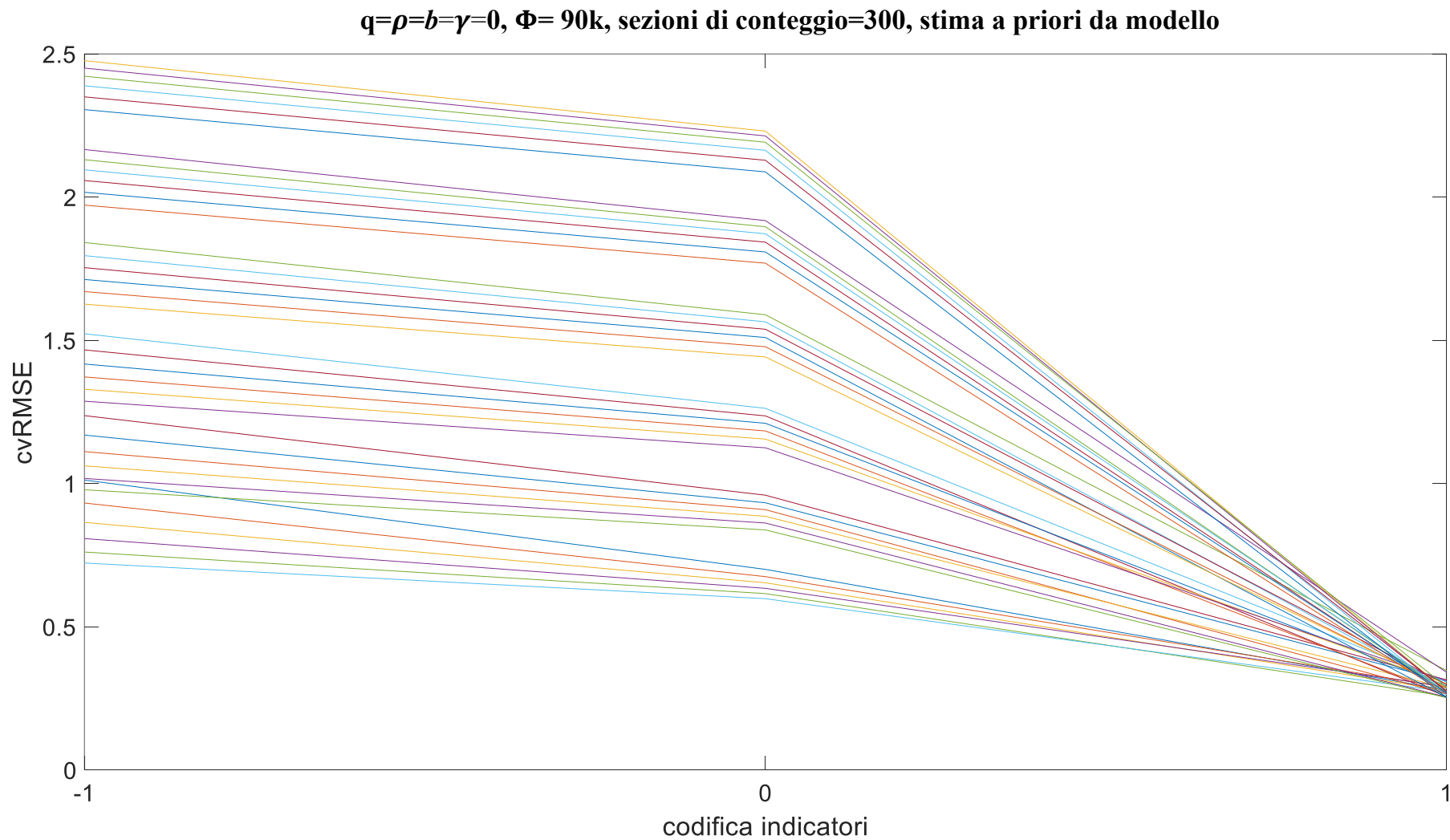


Fig. 5.1 cvRMSE hold out flussi d'arco, $q=\rho=b=\gamma=0$, $\Phi=90k$, sezioni di conteggio=300, stima a priori da modello

La figura 5.1 rappresenta l'andamento dell'indicatore d'errore *cvRMSE* valutato rispetto all'hold out dei flussi d'arco nel caso di $q=\rho=b=\gamma=0$, con tutte le coppie *od* variabili ($\Phi=90$ mil) e 300 sezioni di conteggio. Le diverse curve sono riferite a 36 stima a priori della domanda, ottenute attraverso l'applicazione di un modello sequenziale. Tutte le curve, a prescindere dalla tipologia di distorsione introdotta, hanno andamento decrescente con l'ottimo in corrispondenza del punto (1). Il valore assunto dall'indicatore d'errore *cvRMSE* in corrispondenza del punto iniziale ha una variabilità tra 0.85 e 2.50 (punto -1). A seguito dell'applicazione della procedura di correzione della domanda GLS, lo stesso errore ha una variabilità più stretta e precisamente tra 0.70 e 2.25 (punto 0). Infine, nel punto (1), che mostra i risultati ottenuti applicando la procedura proposta, si osserva un *cvRMSE* con variabilità estremamente bassa e mediamente pari a 0.25. Dunque, la procedura proposta ha consentito un miglioramento estremamente significativo della qualità dell'output, alta in senso assoluto, miglioramento praticamente insensibile, e quindi estremamente robusto, rispetto all'errore iniziale introdotto dalla stima a priori. In particolare, tale miglioramento è compreso tra l'89% (da 2.25 a 0.25) e il 64% (da 0.7 a 0.25). Da notare che i risultati appena commentati, relativi ai punti (-1) e (0), sono riscontrabili anche nella figura (4.3) del capitolo 4. La cosa interessante è che quest'ultima procedura proposta, rispetto a quella riportata nel capitolo 4, consente ulteriori importanti margini di miglioramento (da un *cvRMSE* compreso tra 0.55 e 0.99 a 0.25).

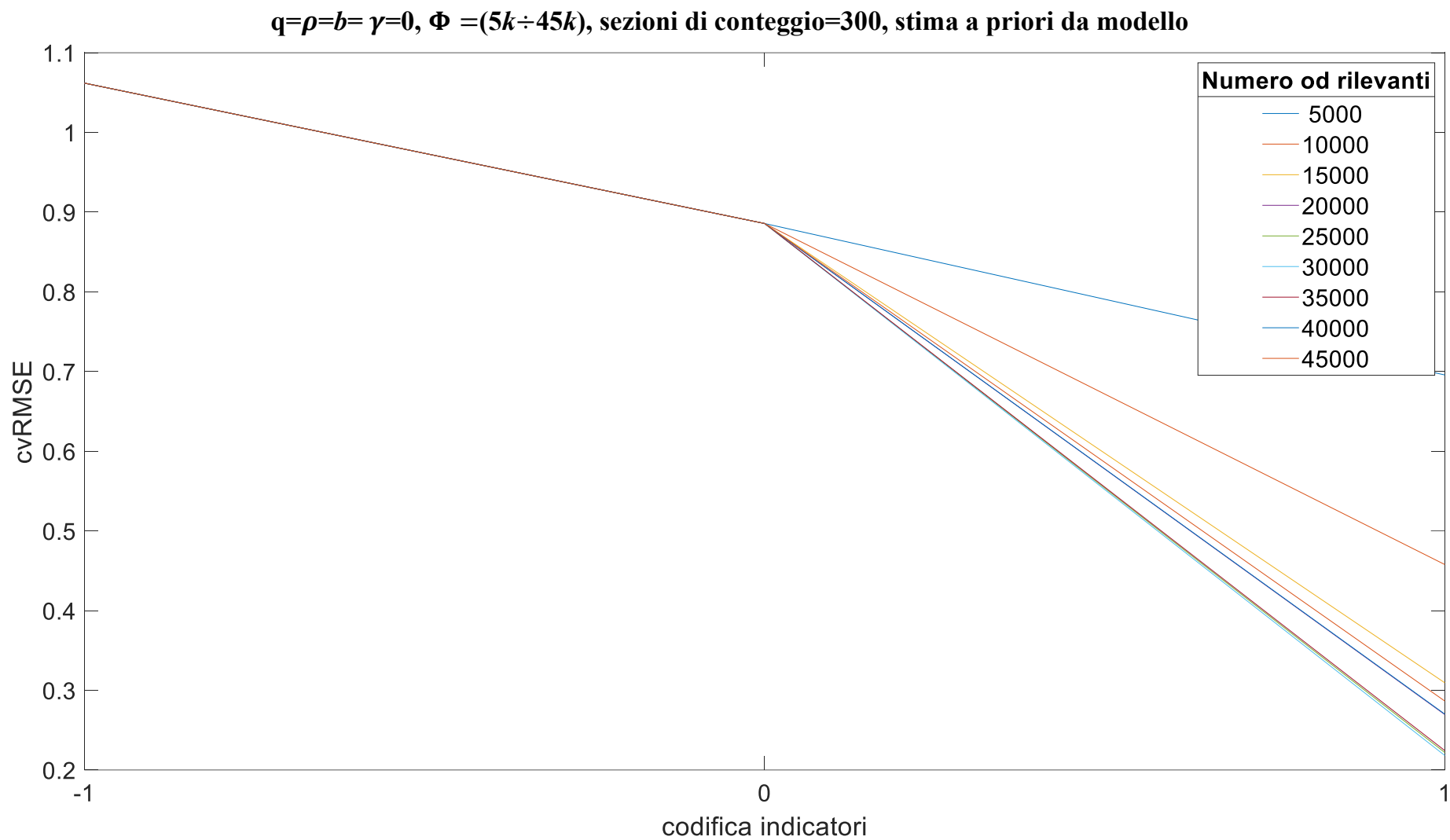


Fig. 5.2 cvRMSE hold out flussi d'arco, $q=\rho=b=\gamma=0$, $\Phi=(5k\div 45k)$, sezioni di conteggio=300, stima a priori da modello

Nella figura 5.2 è rappresentato l'andamento dell'indicatore d'errore $cvRMSE$ rispetto ai flussi d'arco nelle sperimentazioni in cui si è fatto variare Φ , avendo fissato $q=\rho=b=0$, 300 sezioni di conteggio e una stima a priori della domanda ottenuta mediante applicazione di un modello di domanda.

La sperimentazione, quindi, è la stessa del caso precedente, facendo variare però parametricamente il numero Φ delle coppie *od* rilevanti (5k-45k). Tale sperimentazione, per motivi di brevità, è stata effettuata solamente su una delle 36 stime a priori della domanda considerate nella sperimentazione di figura (5.1), visti i risultati analoghi che si sono ottenuti per ognuna di esse.

Con questa tipologia di setting della sperimentazione il numero di incognite presenti in $\mathbf{x}$ è sempre lo stesso, ma tramite esse è necessario riprodurre un minore numero di flussi *od* associati alle Φ coppie *od* rilevanti considerate.

Quando Φ vale 45mila, il valore d'errore letto nel punto (1) è prossimo a 0.25, che è praticamente identico al risultato letto nella figura (5.1) nel punto (1), per i motivi spiegati precedentemente (circa 45mila flussi nulli all'interno dei 90mila elementi del vettore di domanda $\mathbf{d}^{True}$).

È interessante notare che il valore di Φ di miglior compromesso, a cui si accennava precedentemente, si attesta intorno a 30mila, valore per il quale l'indicatore d'errore $cvRMSE$ assume un valore di 0.22.

$q=p=b=\gamma=0$, $\Phi=90k$, sezioni di conteggio= $50 \div 700$, stima a priori da modello

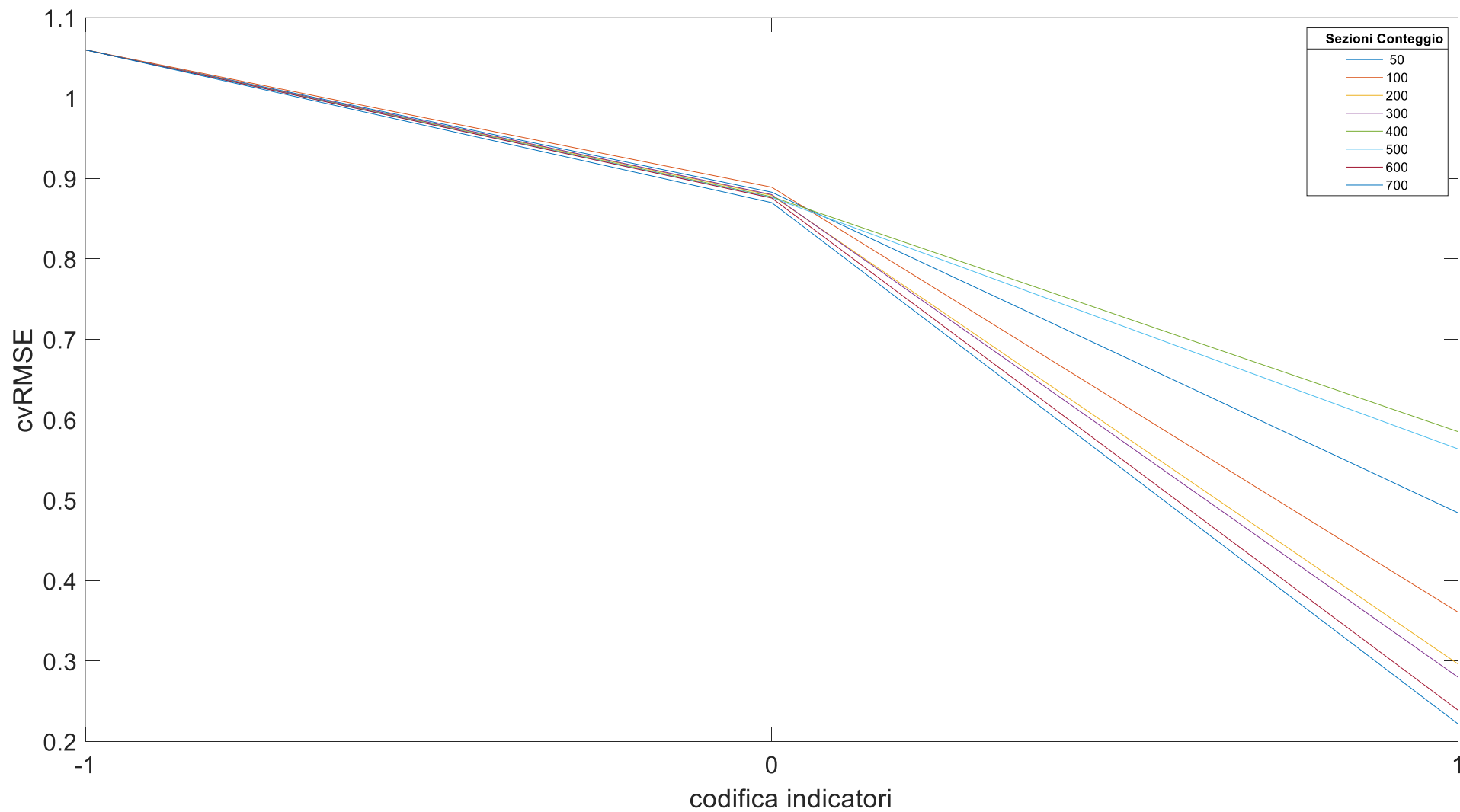


Fig. 5.3 *cvRMSE hold out flussi d'arco, $q=p=b=\gamma=0$, $\Phi = 90k$, sezioni di conteggio= $50 \div 700$, stima a priori da modello*

Nella figura 5.3 è rappresentato l'andamento dell'indicatore d'errore $cvRMSE$ rispetto ai flussi d'arco nelle sperimentazioni in cui si è fatto variare il numero di sezioni di conteggio da 50 a 700, avendo preventivamente fissato $\Phi=90\text{mila}$, $q=\rho=b=0$ e una stima a priori della domanda ottenuta mediante applicazione di un modello di domanda. La domanda utilizzata per eseguire le sperimentazioni è la stessa utilizzata nella sperimentazione precedente (figura 5.2) e cioè una delle 36 utilizzate nella sperimentazione di figura (5.1)

In corrispondenza del punto (0) si verifica che il numero di equazioni disponibili (al più 700) risulta essere notevolmente inferiore rispetto al numero di incognite (circa 90mila). In questi casi il problema risulta notevolmente indeterminato, infatti, qualunque configurazione di conteggi risulta essere insensibile ai fini della correzione della domanda. Al contrario, in corrispondenza del punto (1), il numero di incognite diminuisce e risulta essere pari a $m=2N$ (604), mentre il numero di equazioni varia tra 50 e 700. In questo caso il numero e la localizzazione dei conteggi hanno molta rilevanza, facendo variare in modo significativo il bilancio tra equazioni ed incognite: come risultato, l'indicatore d'errore $cvRMSE$ varia tra 0.22 (con 700 conteggi) e 0.48 (con 50 conteggi).

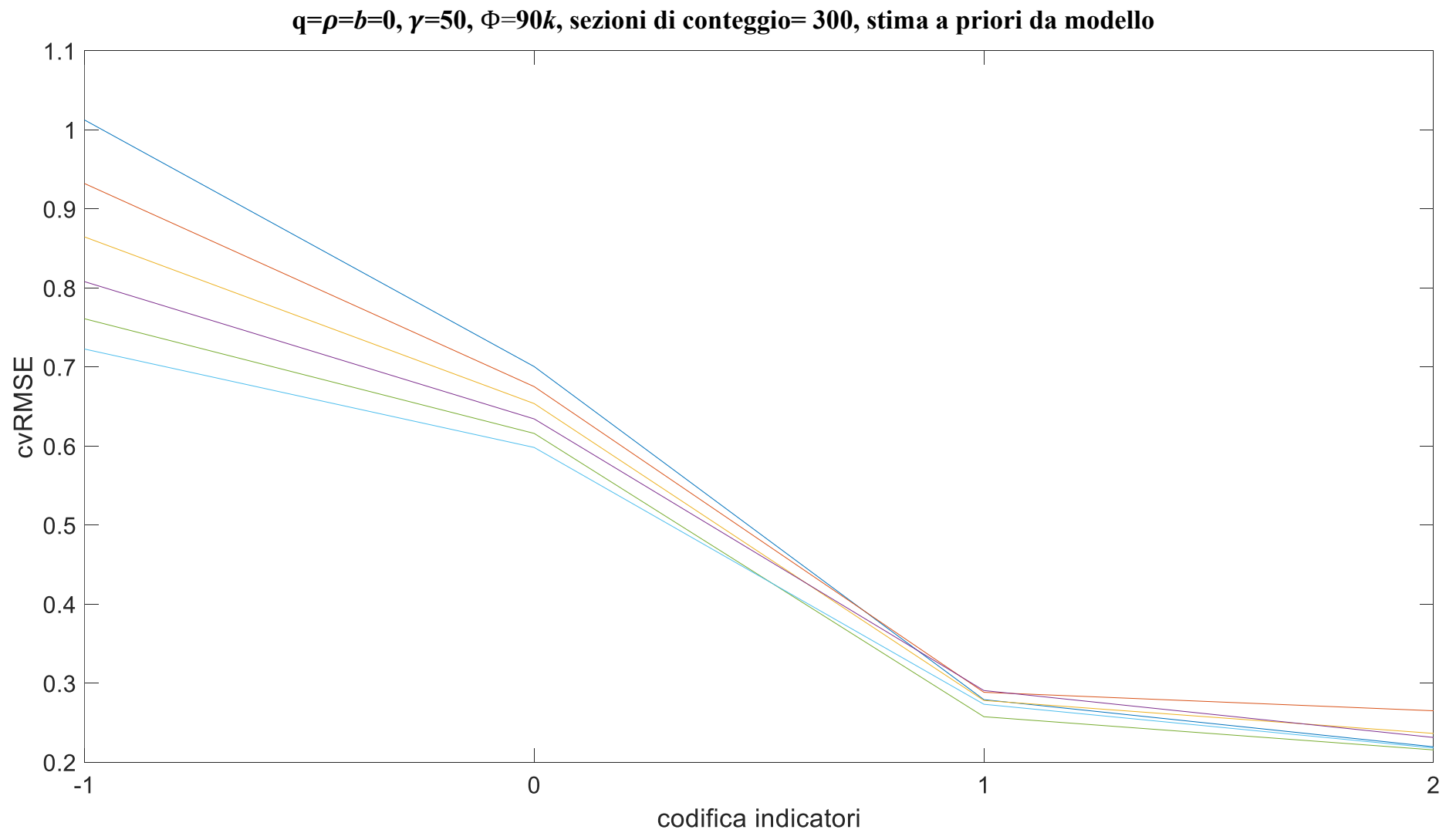


Fig. 5.4 cvRMSE hold out flussi d'arco, $q=\rho=b=0, \gamma=50, \Phi=90k$, sezioni di conteggio: 300, stima a priori da modello

I risultati riportati nella figura (5.4) rappresentano l'andamento dell'indicatore $cvRMSE$ valutato rispetto all'hold out dei flussi d'arco nell'ipotesi di $\gamma=50$, $q=p=b=0$, $\Phi=90k$, 300 sezioni di conteggio e stima a priori della domanda proveniente da modello.

Come già chiarito all'inizio di questo paragrafo, aver settato il parametro γ ad un valore maggiore di zero, oltre all'incremento del numero di incognite da stimare, comporta una modifica nella rappresentazione grafica dei risultati, cioè l'asse orizzontale è costituito da quattro punti e non più tre. I punti (-1) e (0) mantengono lo stesso significato assunto nei grafici precedenti, al punto (1) corrisponde il risultato ottenibile attraverso la procedura proposta assumendo $\gamma=0$, cioè senza l'introduzione di ulteriori incognite, al punto (2) corrisponde il risultato ottenuto a seguito dell'applicazione della procedura proposta con $\gamma \neq 0$, nel caso specifico pari a 50.

La procedura è stata applicata a solo 6 delle 36 stime a priori della domanda che compaiono nella figura (5.1).

Come si evince dal grafico, l'andamento dell'errore $cvRMSE$ decresce, con un minimo proprio in corrispondenza del punto (2), ciò significa che aver fornito al sistema ulteriori gradi di libertà mediante l'introduzione delle ulteriori incognite γ , ha comportato un significativo beneficio rispetto al caso $\gamma=0$ punto (1). In altri termini, la capacità esplicativa delle variabili γ , legate al tempo di percorrenza tra le varie coppie od , sembra più che compensare il corrispondente aumento dello sbilancio tra equazioni ed incognite.

Nelle figure 5.5 sono riportati dei risultati del tutto analoghi a quelli presenti nella figura 5.4. L'unica differenza in termini di input è il parametro γ assunto pari a 100 (invece di 50). Anche in questo caso l'indicatore d'errore $cvRMSE$ ha un andamento decrescente con l'ottimo proprio in corrispondenza dell'ultimo punto a cui corrispondono un numero maggiore di incognite da ottimizzare. Confrontando le due figure 5.5 e 5.4, al di là dell'ovvia coincidenza dei punti (-1,0,1), si osservano delle performance ancora migliori nel caso di $\gamma = 100$: nel punto (2) con $\gamma = 50$ si raggiunge un valore del $cvRMSE$ mediamente superiore a 0.23, invece, con $\gamma = 100$ il $cvRMSE$ assume un valore prossimo a 0.21.

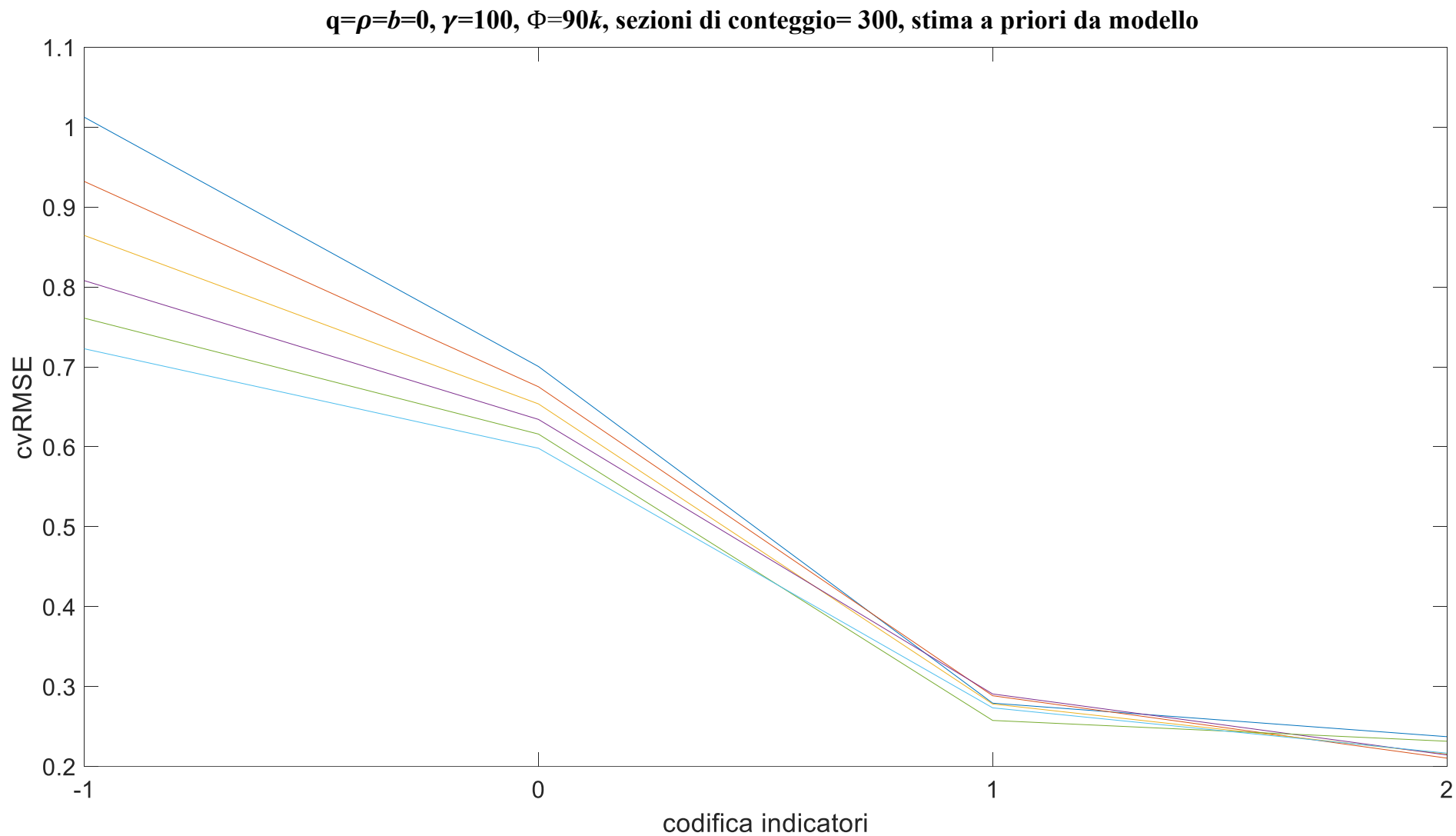


Fig. 5.5 cvRMSE hold out flussi d'arco, $q=\rho=b=0, \gamma=100, \Phi=90k$, sezioni di conteggio= 300, stima a priori da modello

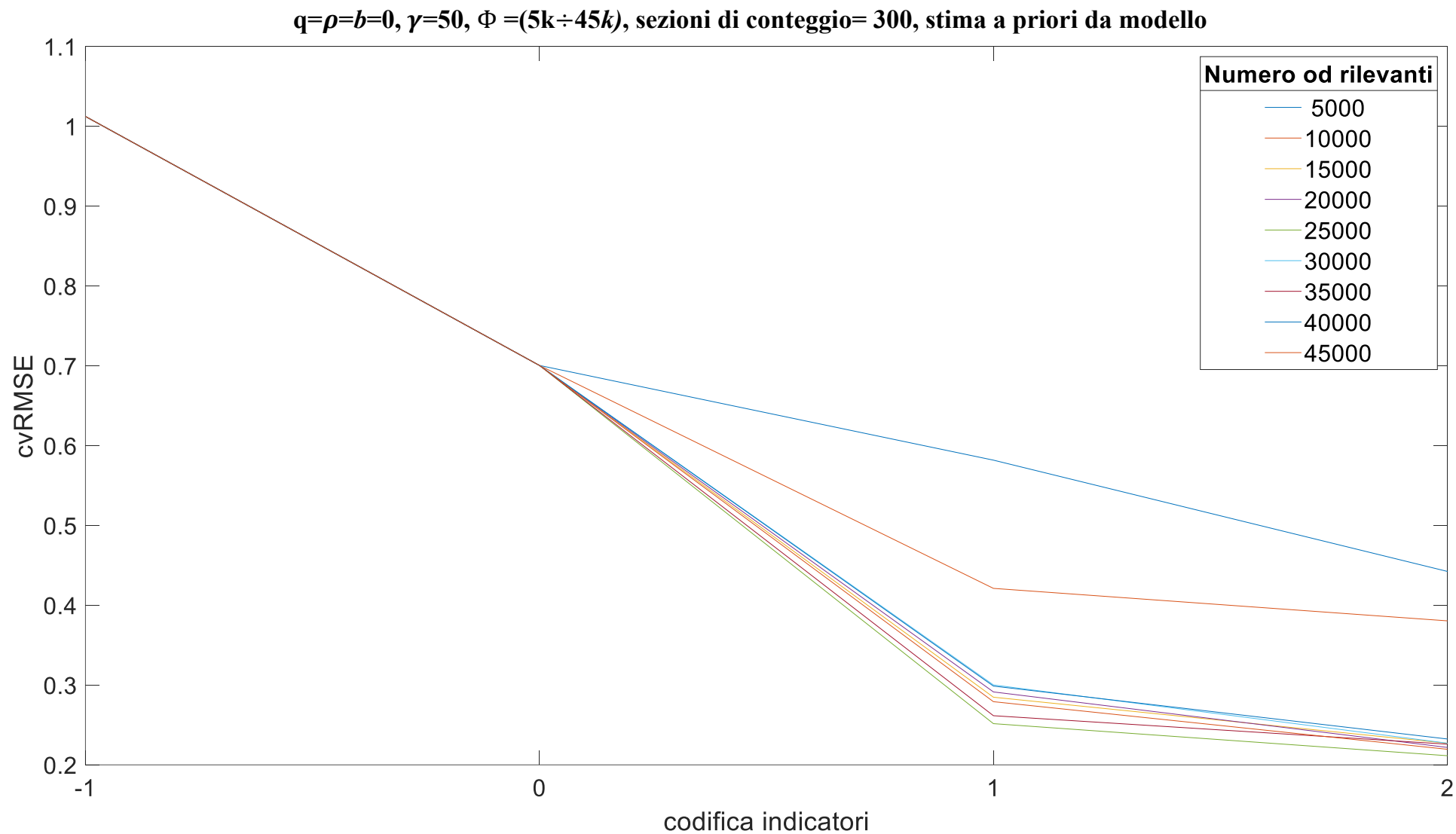


Fig. 5.6 cvRMSE hold out flussi d'arco, $q=\rho=b=0, \gamma=50, \Phi=5k \div 45k$, sezioni di conteggio= 300, stima a priori da modello

Nella figura 5.6 è rappresentato l'andamento dell'indicatore d'errore *cvRMSE* rispetto ai flussi d'arco nelle sperimentazioni in cui si è fatto variare Φ , avendo fissato $\gamma=50$, $q=p=b=0$, 300 sezioni di conteggio e una stima a priori della domanda ottenuta mediante applicazione di un modello di domanda. Il procedimento è del tutto analogo a quello raccontato in figura (5.2), con l'aggiunta dell'ulteriore variabile $\gamma \neq 0$.

In questo caso la procedura è stata applicata ad una sola delle 6 domande riportate nella figura (5.4). I grafici mostrati possono essere intesi come uno studio di dettaglio su alcuni dei risultati ottenuti in figura (5.4), infatti, entrambi gli studi si riferiscono alla condizione $\gamma = 50$, con l'aggiunta della variabilità del nel numero di *od* rilevanti. Eseguendo un'analisi dei risultati si può osservare che il valore minimo dell'indicatore *cvRMSE*, pari a 0.21, lo si ottiene nel punto (2) con Φ pari a 25mila, ed è migliore di quello ottenibile considerando tutte le coppie *od*, che si ricorda essere di poco superiore a 0.23 (cfr. figura 5.4).

I risultati rappresentati nella figura (5.7) sono del tutto analoghi a quelli presentati nella figura (5.6), pertanto, tutte le considerazioni fatte precedentemente possono essere ripetute. L'unica differenza tra le due figure riguarda il parametro di input γ , che in questo caso assume il valore 100 (invece di 50). Per quanto riguarda il valore minimo dell'indicatore *cvRMSE* (pari a 0.20), si realizza sempre nel punto (2) per un valore di Φ paria a 20mila e, anche in questo caso, è migliore di quello ottenibile considerando tutte le coppie *od*, che si ricorda essere di circa 0.21 (cfr. figura 5.4). E' importante sottolineare come, quando si raggiungono valori bassi in senso assoluto del *cvRMSE*, ottenere miglioramenti ulteriori diventa sempre più difficile e, quindi, anche piccoli miglioramenti sono da considerarsi importanti.

$q=\rho=b=0, \gamma=100, \Phi=(5k \div 45k)$, sezioni di conteggio: 300, stima a priori da modello

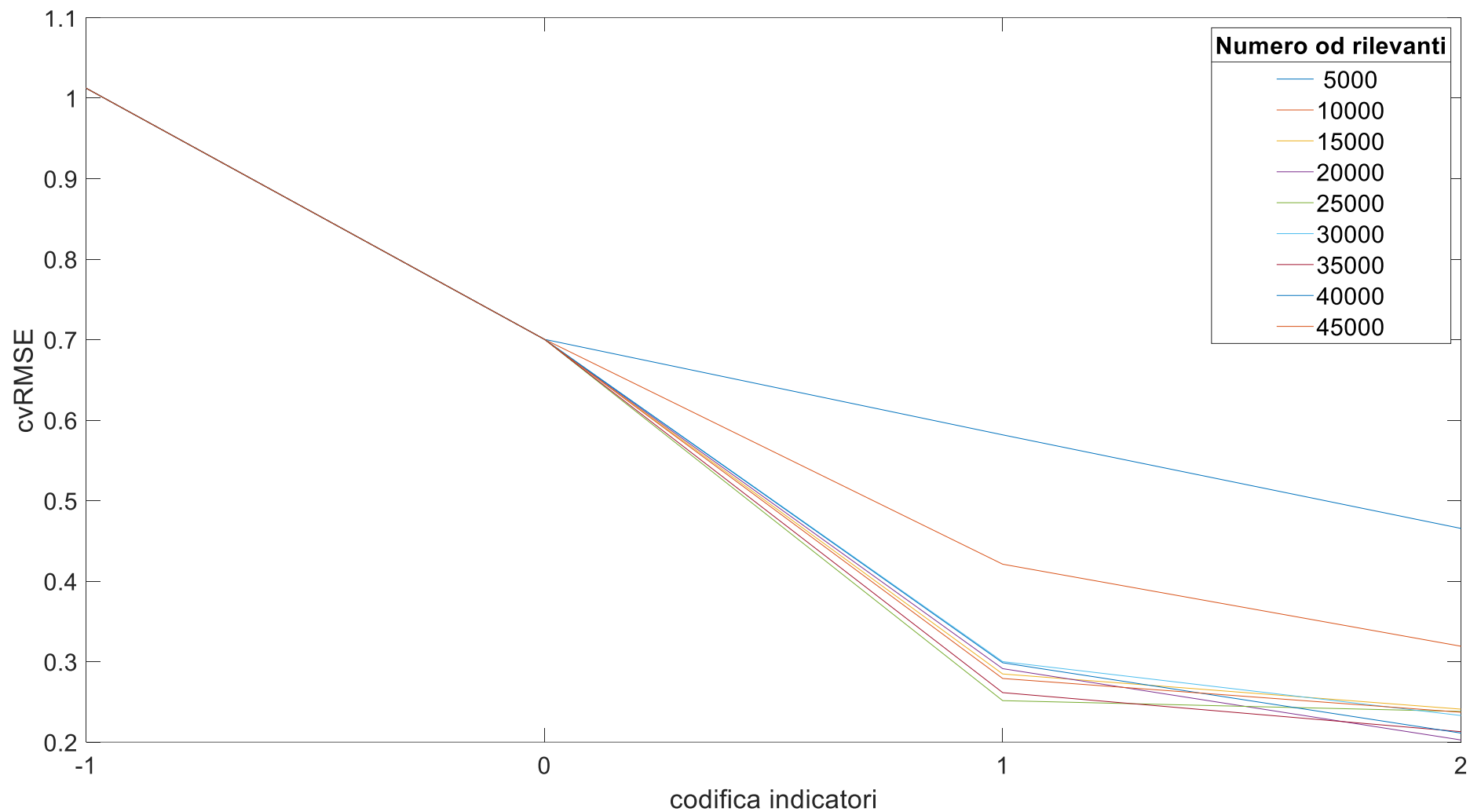


Fig. 5.7 *cvRMSE hold out flussi d'arco, $q=\rho=b=0, \gamma=100, \Phi=(5k \div 45k)$, sezioni di conteggio= 300, stima a priori da modello*

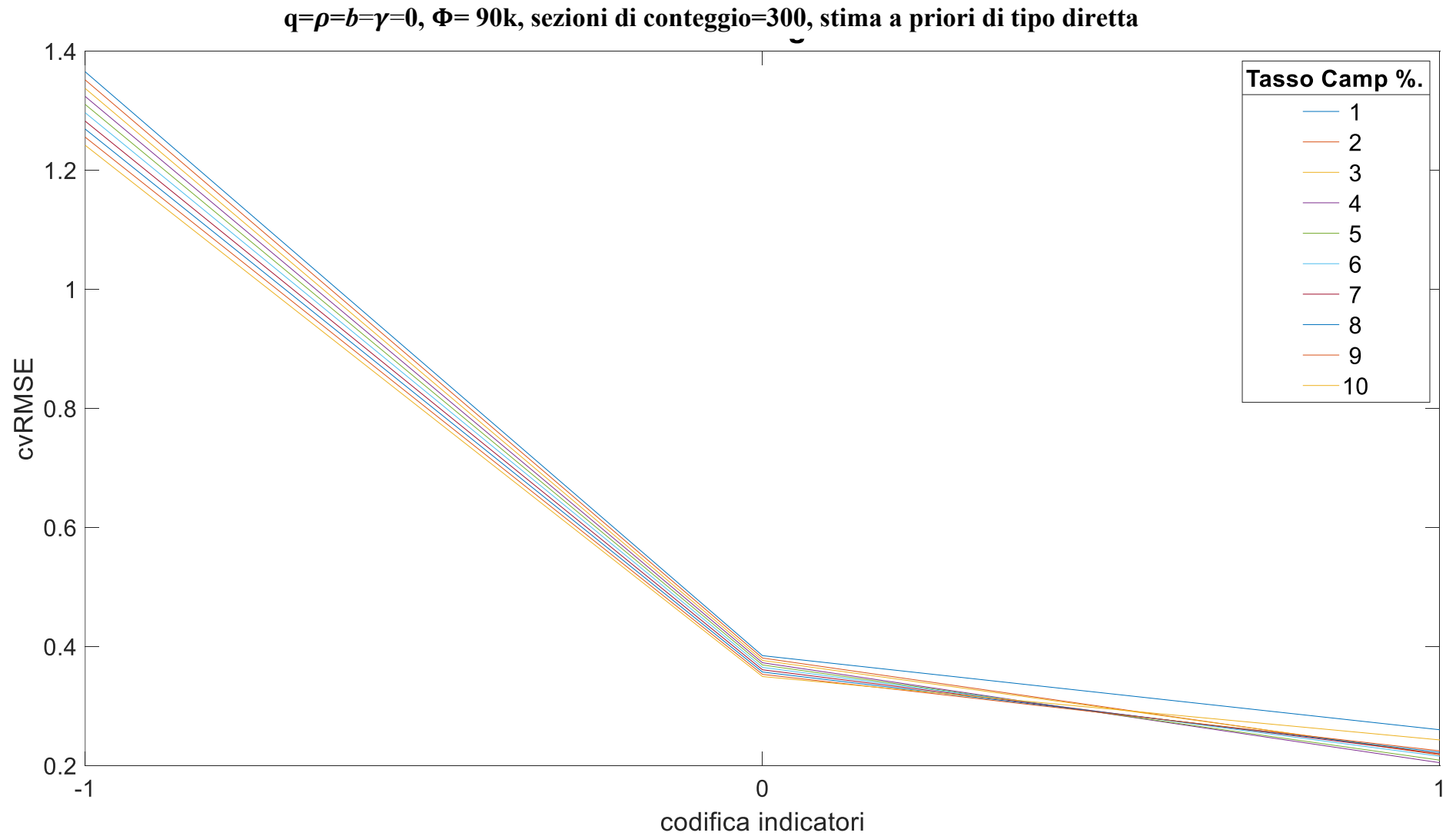


Fig. 5.8 *cvRMSE hold out flussi d'arco, $q=\rho=b=\gamma=0$, $\Phi=90k$, sezioni di conteggio=300, stima a priori di tipo diretta*

La figura 5.8 rappresenta l'andamento dell'indicatore d'errore $cvRMSE$ valutato rispetto all'hold out dei flussi d'arco nel caso di $q=\rho=b=\gamma=0$, con $\Phi=90$ mila e 300 sezioni di conteggio. Le diverse curve sono riferite a 10 stime a priori della domanda, ottenute attraverso l'applicazione di una procedura di stima diretta della domanda per altrettanti 10 tassi di campionamento, variabili tra 1% e 10%.

È interessante notare che i risultati sono pressoché insensibili al tasso di campionamento, ciononostante la procedura proposta consente il raggiungimento di un errore sui flussi d'arco medio prossimo a 0.2 (punto 1). L'omologo valore ottenibile attraverso l'applicazione della procedura classica di correzione GLS della domanda è prossimo a 0.38 (punto 0), dunque si osserva un miglioramento del 47% (da 0.38 a 0.2). Il miglioramento inferiore in termini relativi che si osserva dall'applicazione della procedura proposta rispetto a quella classica, in questa sperimentazione nella quale la stima a priori della domanda è stata ottenuta con il metodo della stima diretta invece che con la stima da modello, è dovuto al fatto che con tale stima a priori anche la correzione GLS classica raggiunge dei valori del $cvRMSE$ piuttosto bassi e quindi, come già ricordato in precedenza, è più difficile ottenere miglioramenti ulteriori.

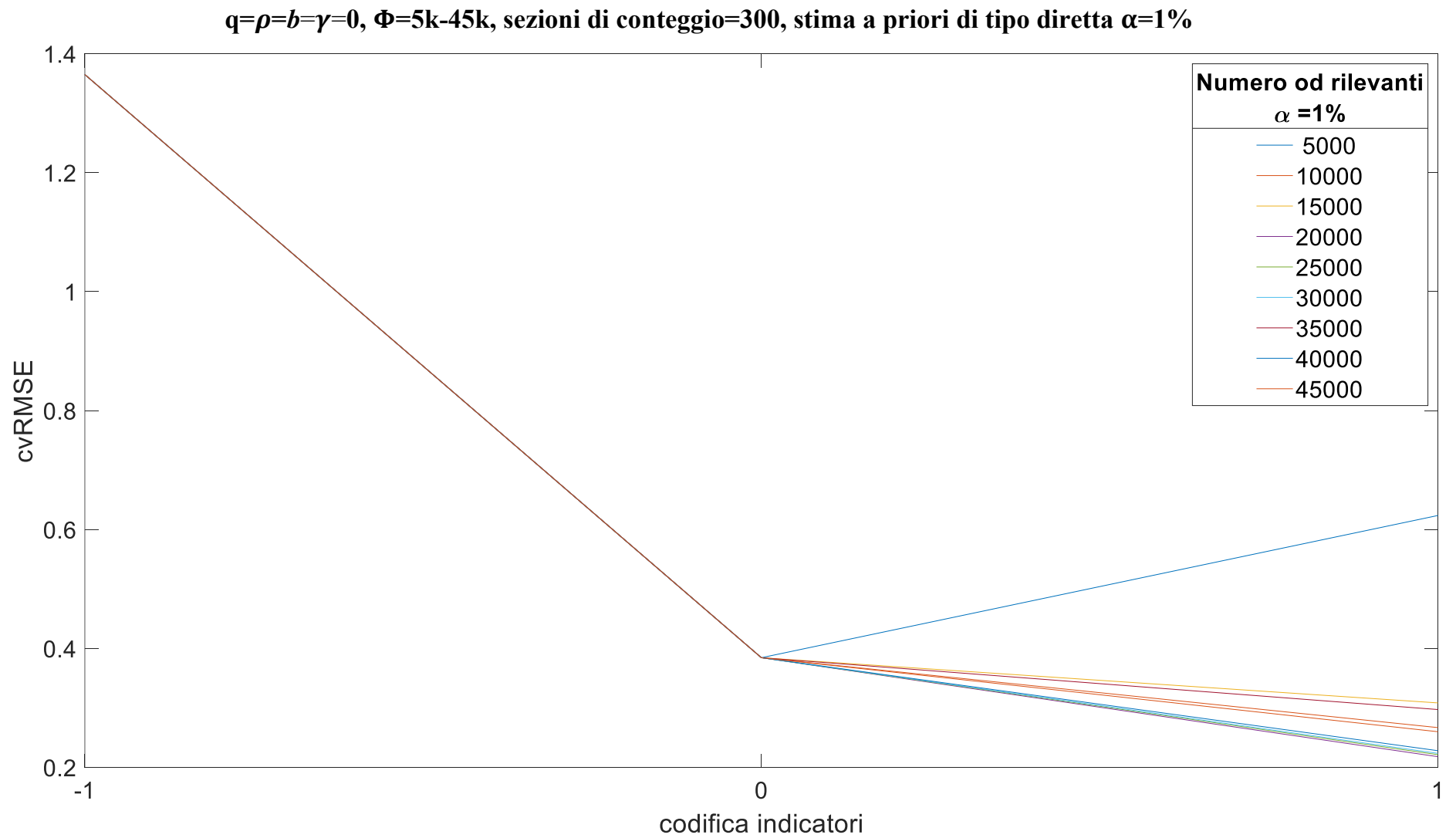


Fig. 5.9 *cvRMSE hold out flussi d'arco, $q=\rho=b=\gamma=0$, $\Phi = 5k-45k$, sezioni di conteggio=300, stima a priori di tipo diretta $\alpha=1\%$*

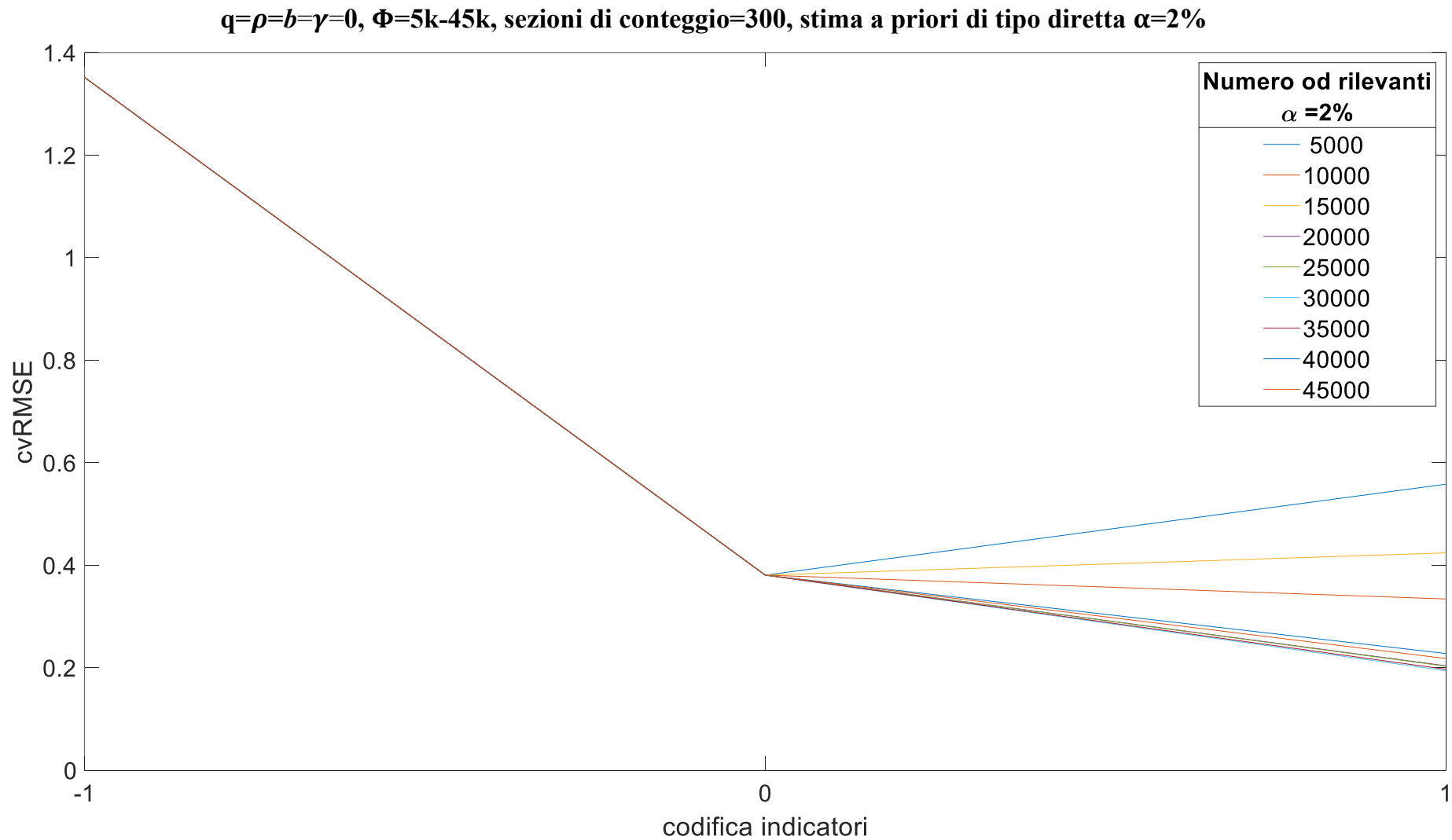


Fig. 5.10 cvRMSE hold out flussi d'arco, $q=\rho=b=\gamma=0$, $\Phi = 5k-45k$, sezioni di conteggio=300, stima a priori di tipo diretta $\alpha=2\%$

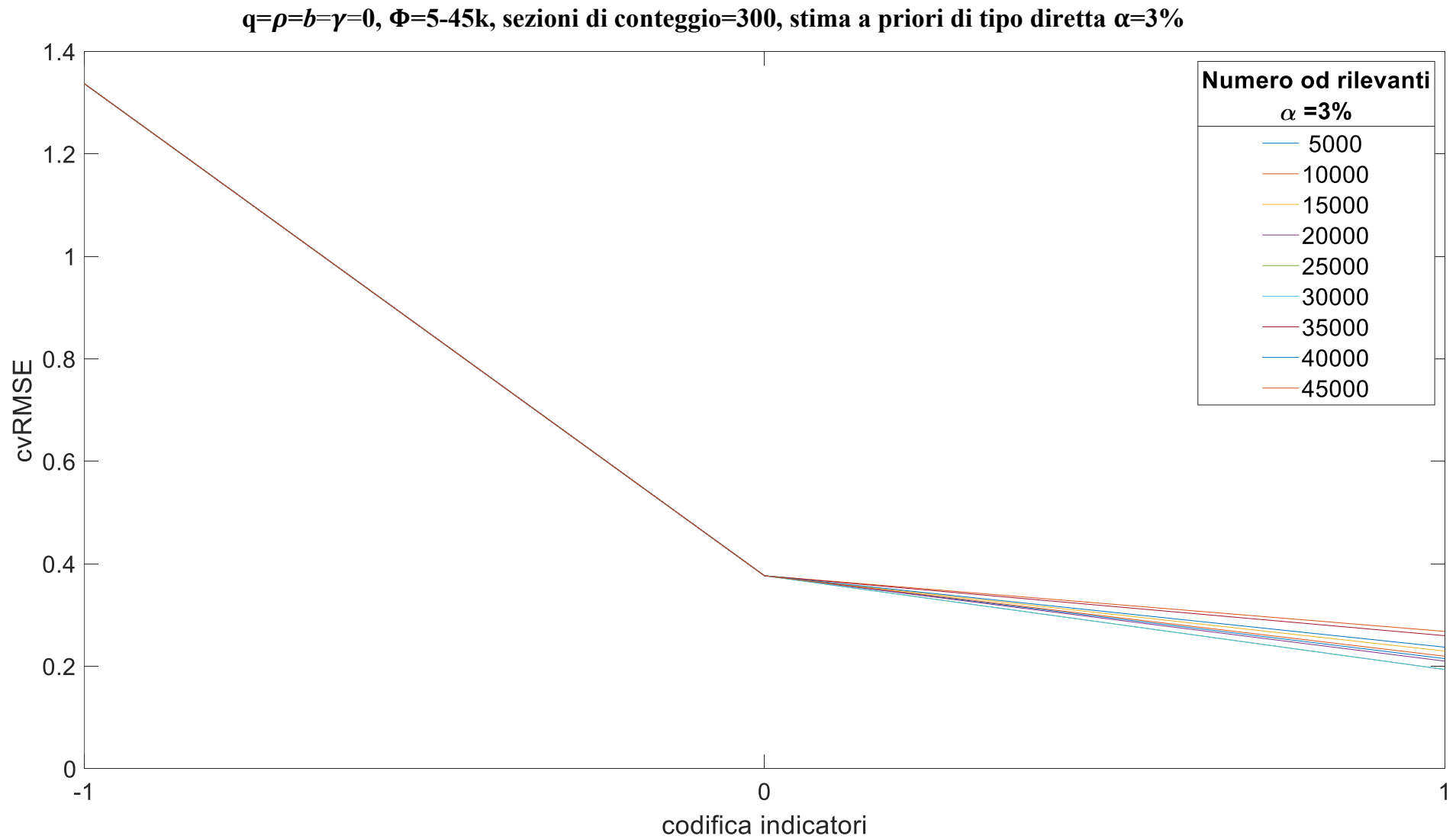


Fig. 5.11 cvRMSE hold out flussi d'arco, $q=\rho=b=\gamma=0$, $\Phi = 5k-45k$, sezioni di conteggio=300, stima a priori di tipo diretta $\alpha=3\%$

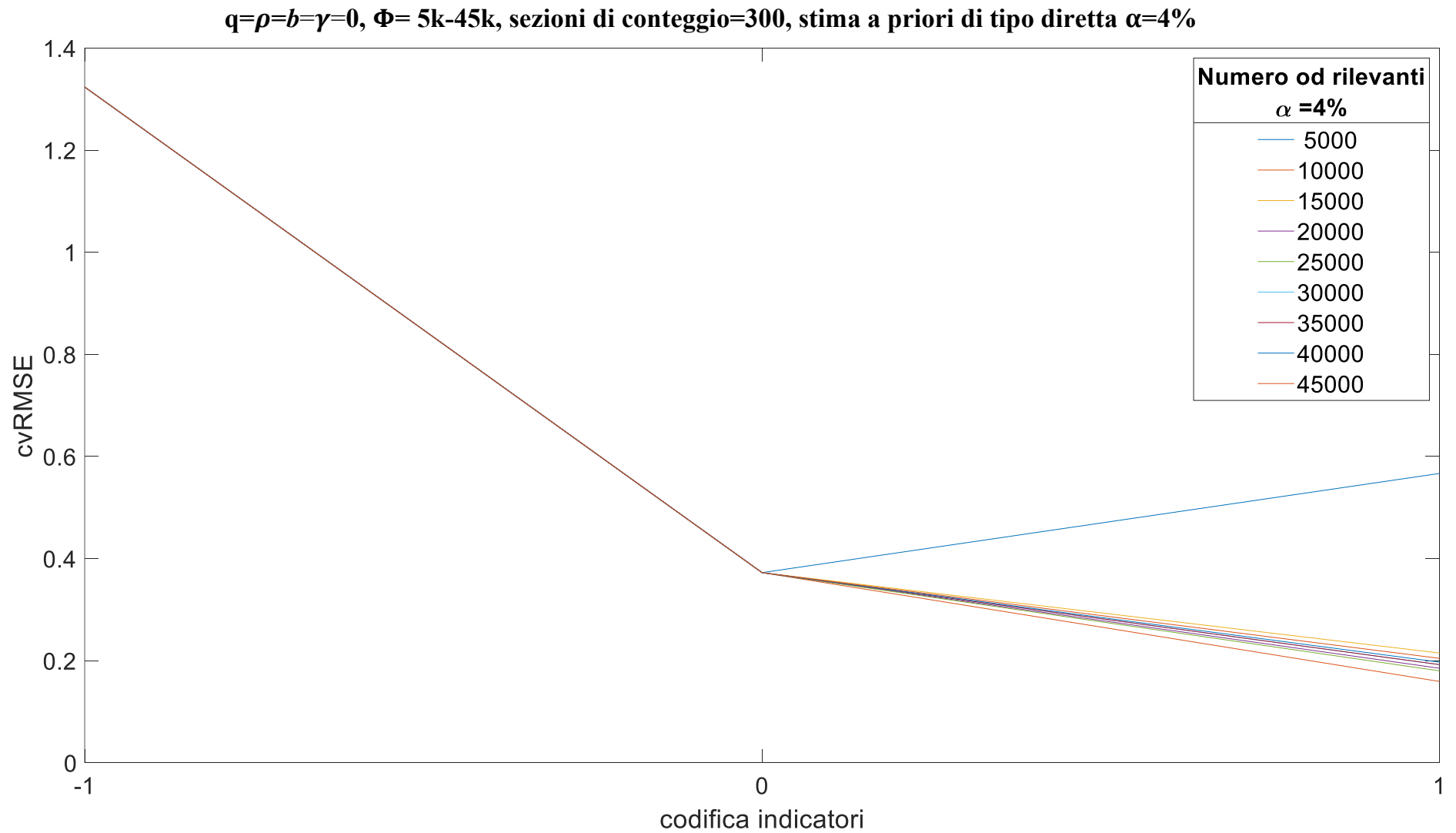


Fig. 5.12 cvRMSE hold out flussi d'arco, $q=\rho=b=\gamma=0$, $\Phi =5k-45k$, sezioni di conteggio=300, stima a priori di tipo diretta $\alpha=4\%$

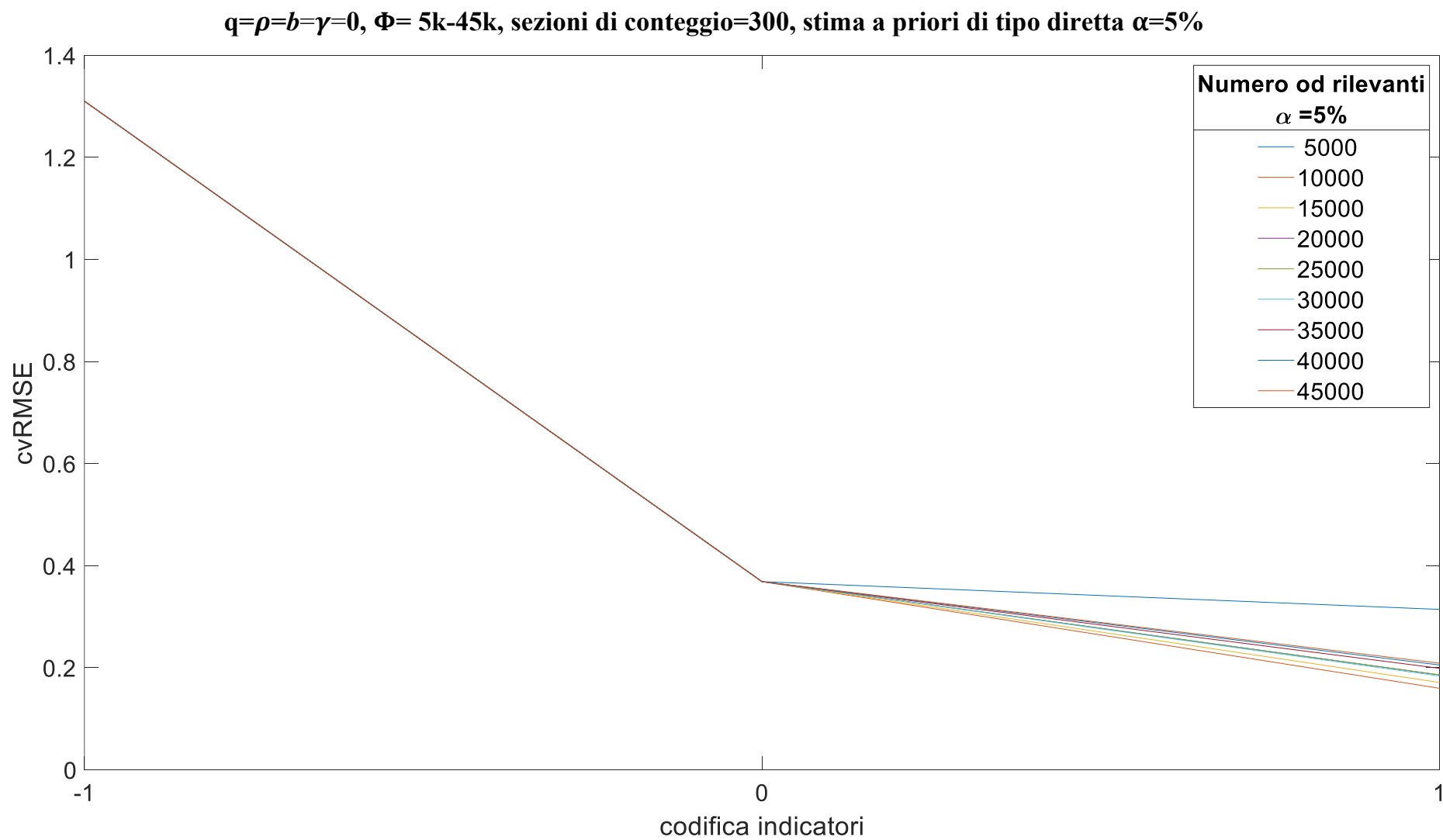


Fig. 5.13 *cvRMSE hold out flussi d'arco*, $q=\rho=b=\gamma=0$, $\Phi=4k-45k$, sezioni di conteggio=300, stima a priori di tipo diretta $\alpha=5\%$

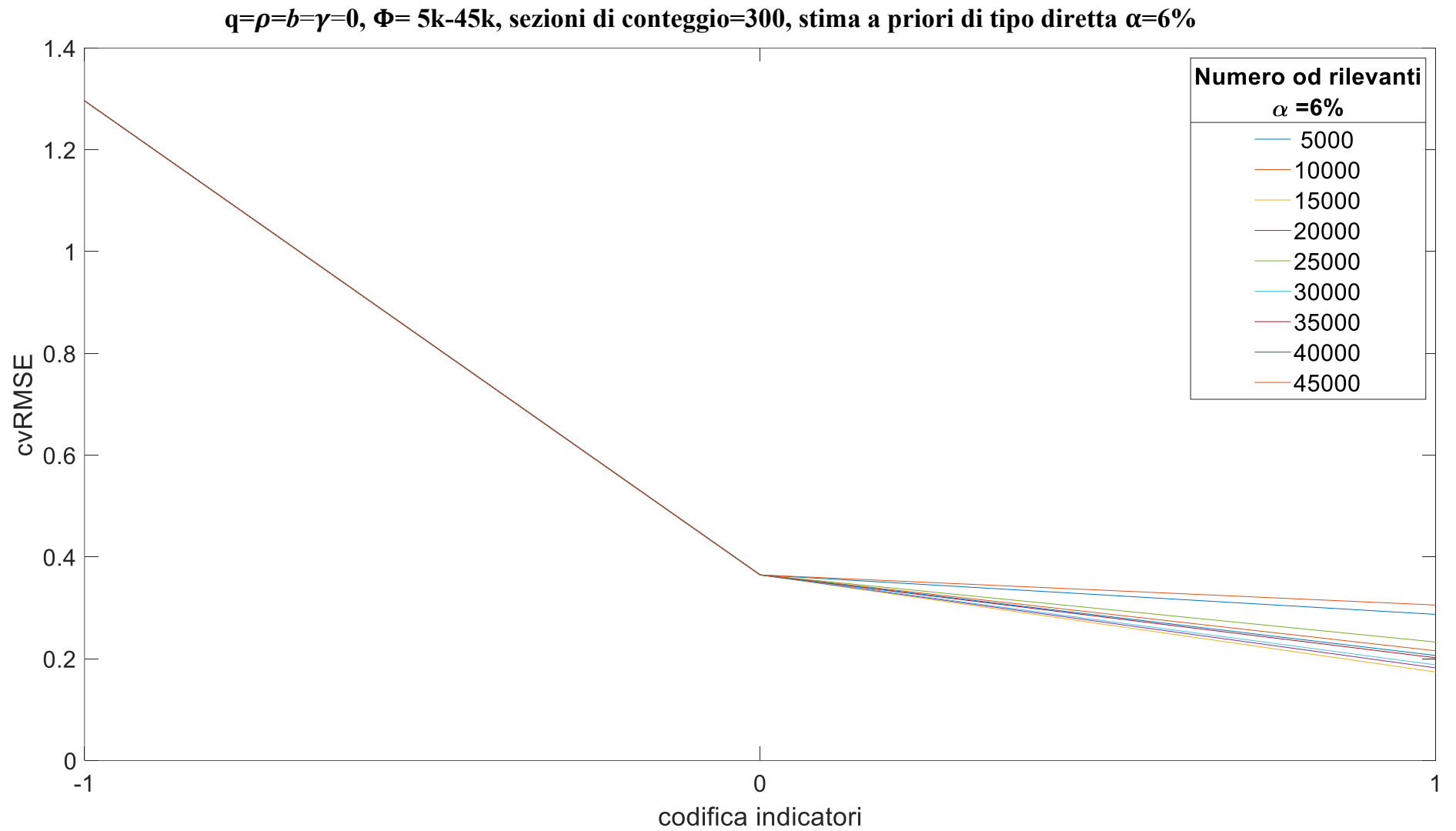


Fig. 5.14 *cvRMSE hold out flussi d'arco*, $q=\rho=b=\gamma=0$, $\Phi =5k-45k$, sezioni di conteggio=300, stima a priori di tipo diretta $\alpha=6\%$

$q=\rho=b=\gamma=0$, $\Phi=5k-45k$, sezioni di conteggio=300, stima a priori di tipo diretta $\alpha=7\%$

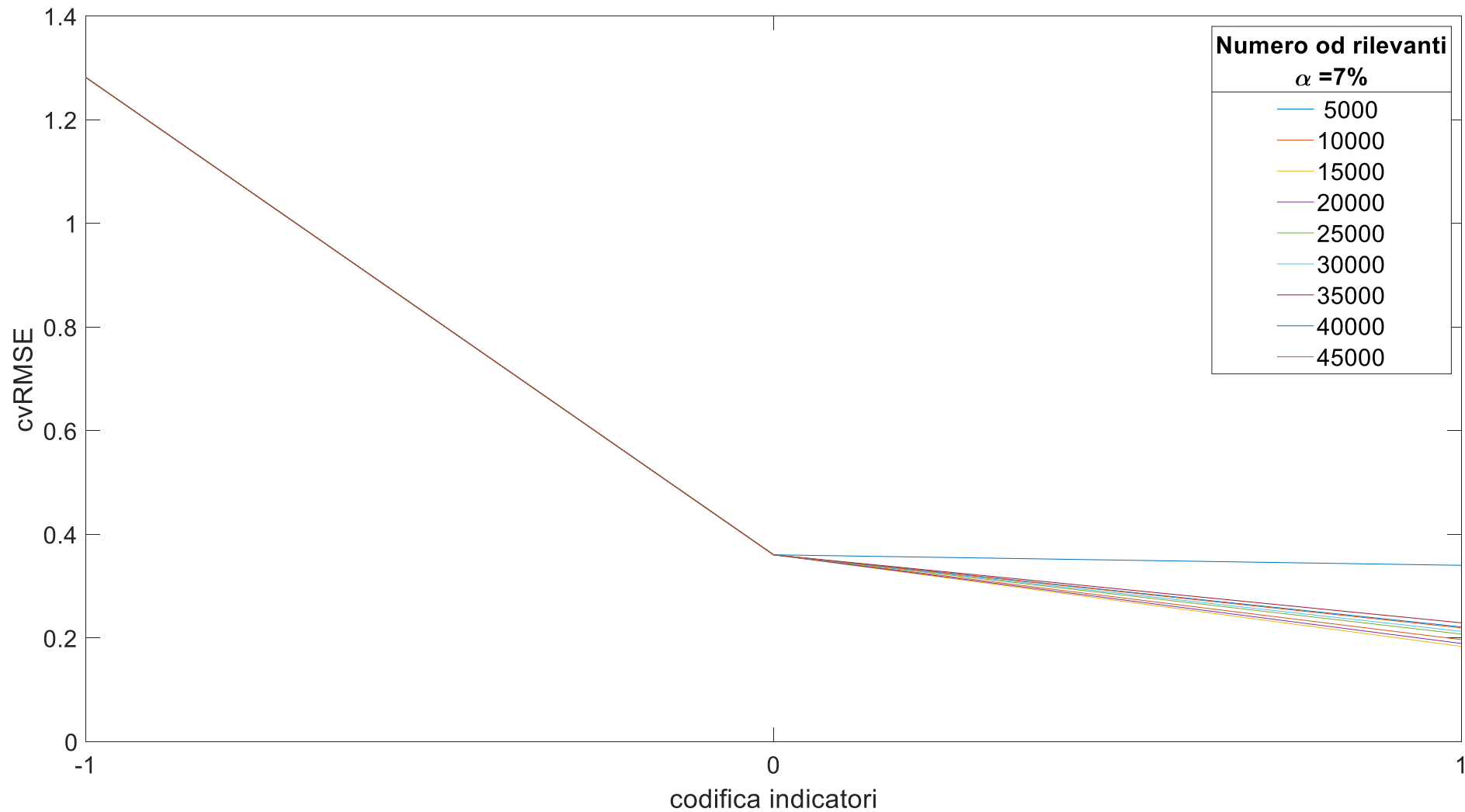


Fig. 5.15 cvRMSE hold out flussi d'arco, $q=\rho=b=\gamma=0$, $\Phi=5k-45k$, sezioni di conteggio=300, stima a priori di tipo diretta $\alpha=7\%$

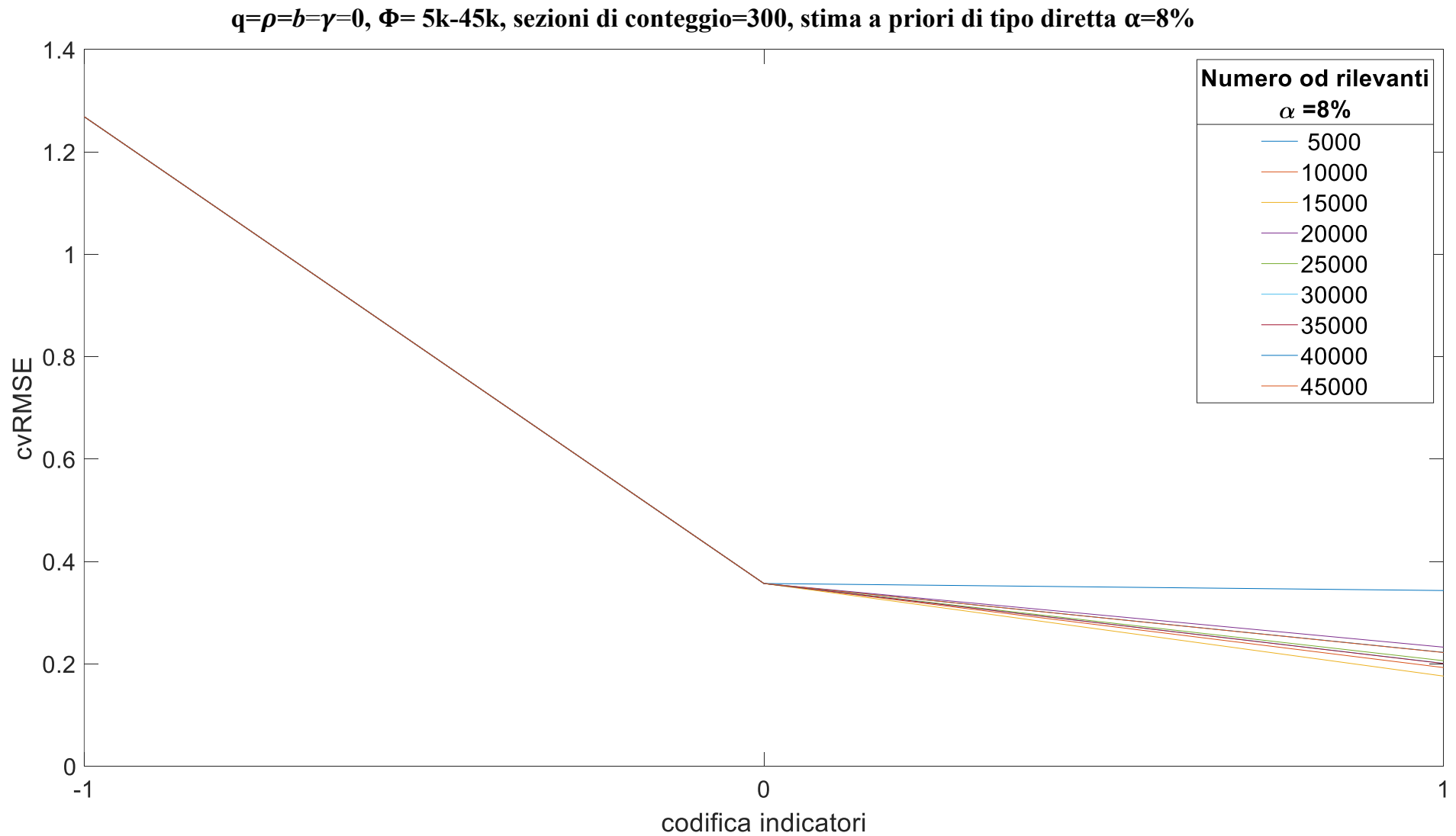


Fig. 5.16 cvRMSE hold out flussi d'arco, $q=\rho=b=\gamma=0$, $\Phi=5k-45k$, sezioni di conteggio=300, stima a priori di tipo diretta $\alpha=8\%$

$q=\rho=b=\gamma=0$, $\Phi= 5\text{-}45\text{k}$, sezioni di conteggio=300, stima a priori di tipo diretta $\alpha=9\%$

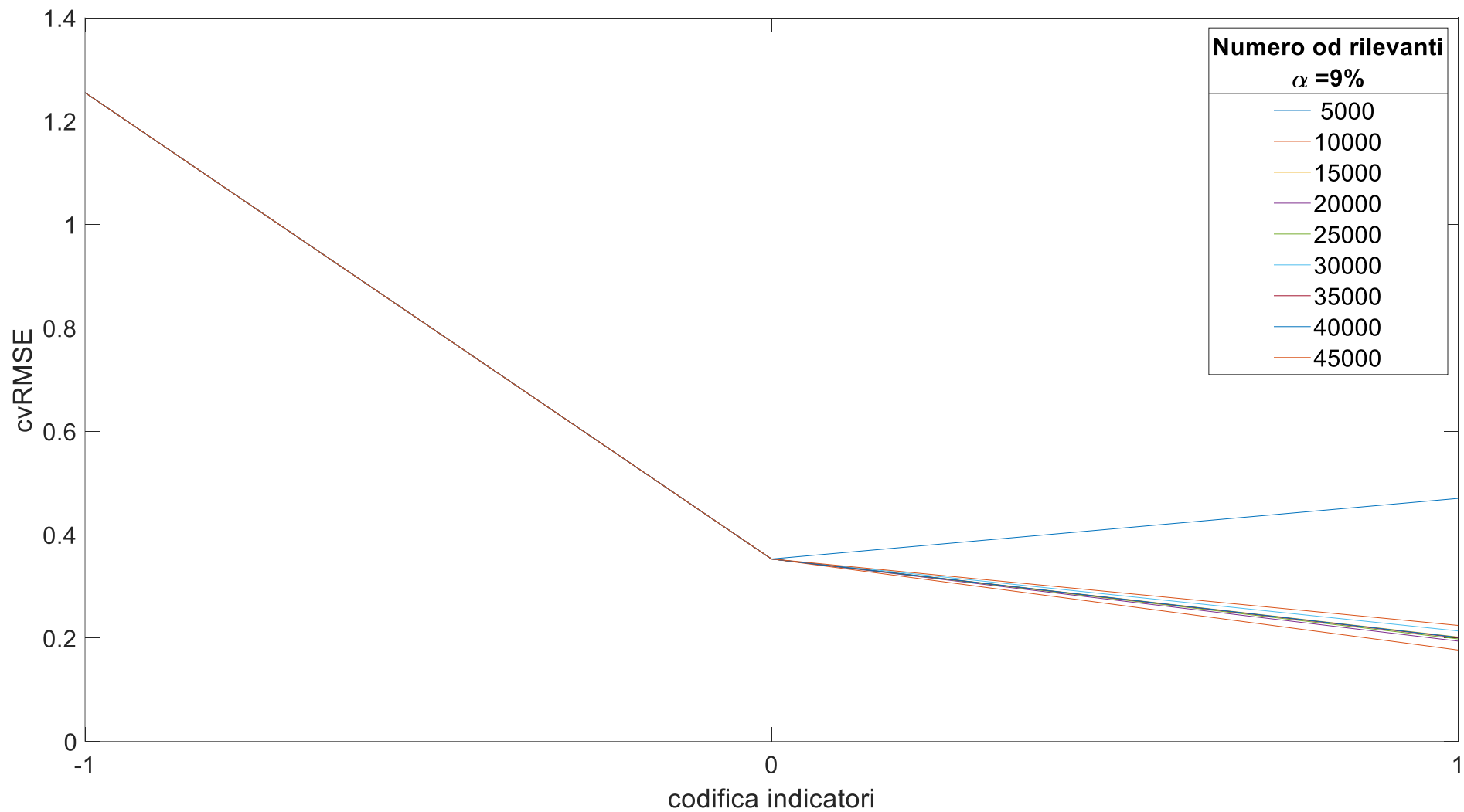


Fig. 5.17 *cvRMSE hold out flussi d'arco*, $q=\rho=b=\gamma=0$, $\Phi = 5\text{k-}45\text{k}$, sezioni di conteggio=300, stima a priori di tipo diretta $\alpha=9\%$

$q=\rho=b=\gamma=0$, $\Phi= 5k-45k$, sezioni di conteggio=300, stima a priori di tipo diretta $\alpha=10\%$

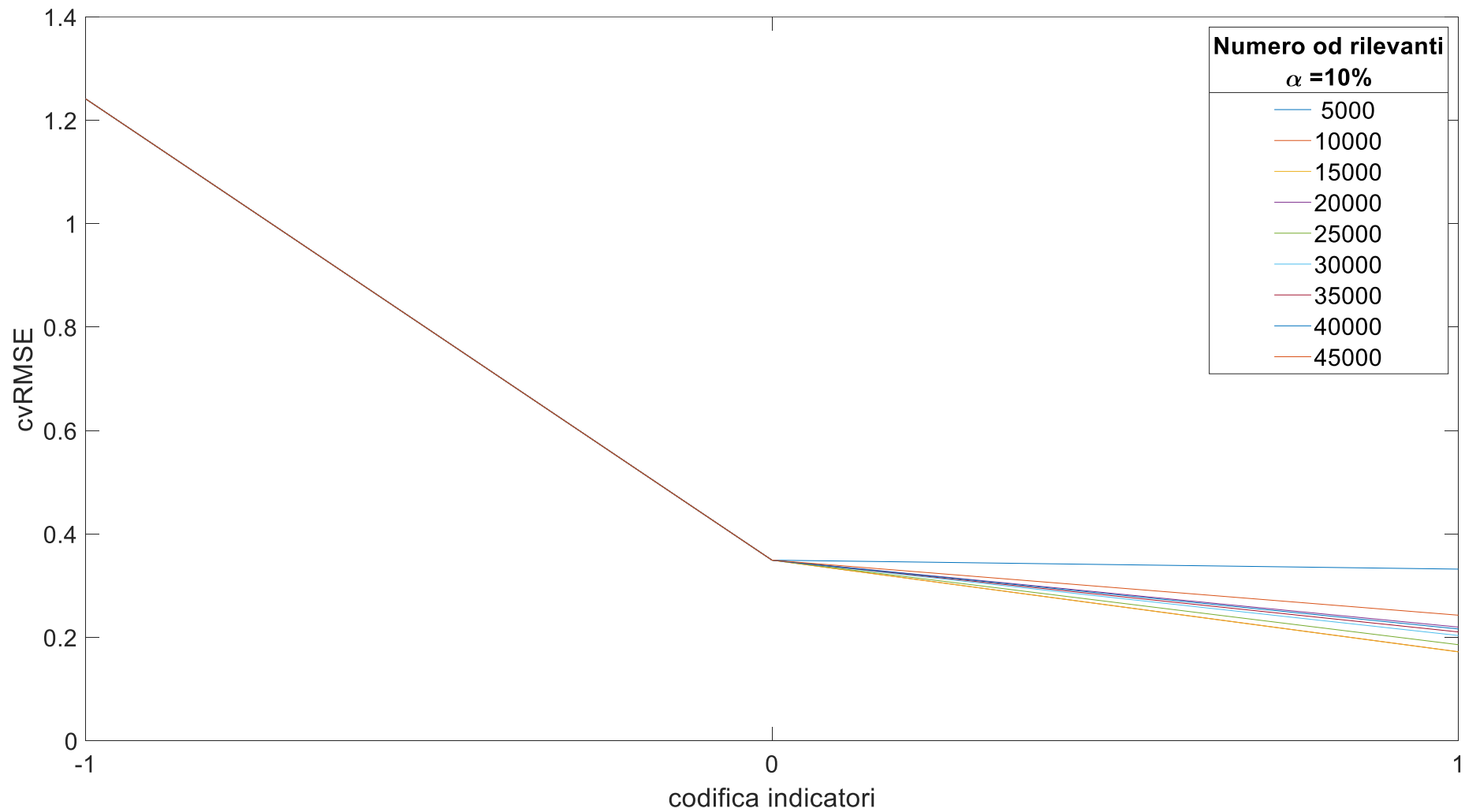


Fig. 5.18 *cvRMSE hold out flussi d'arco*, $q=\rho=b=\gamma=0$, $\Phi = 5k-45k$, sezioni di conteggio=300, stima a priori di tipo diretta $\alpha=10\%$

I risultati nelle figure dalla (5.9) alla (5.17) rappresentano l'andamento dell'indicatore d'errore $cvRMSE$ valutato rispetto all'hold out dei flussi d'arco nel caso di $q=\rho=b=\gamma=0$ e 300 sezioni di conteggio. Tra le diverse figura varia la stima a priori della domanda, ottenuta mediante una metodologia di stima diretta per un prefissato tasso di campionamento $\alpha \in [1\% \div 10\%]$. All'interno della stessa figura è stata studiata l'efficacia della procedura di correzione proposta al variare di Φ , analogamente a quanto già fatto nelle figure (5.2, 5.6 e 5.7).

Al crescere del tasso di campionamento si riduce sia l'errore iniziale (punto -1) si passa da 1.34 ($\alpha=1\%$) a 1.26 ($\alpha=10\%$), sia l'errore dovuto alla procedura di correzione classica (punto 1) da 0.38 ($\alpha=1\%$) a 0.35 ($\alpha=10\%$). Nel punto (1), invece, è necessario considerare oltre alla variabilità di α anche quella di Φ . Infatti, per $\alpha \leq 3\%$ il valore più basso dell'indicatore si realizza per Φ compreso tra 20mila e 30mila, negli altri casi ($\alpha > 3\%$) il minimo si realizza per Φ compreso tra 10mila e 15 mila. In generale il valore più basso della procedura di correzione è 0.17 in corrispondenza di $\alpha=10\%$ e $\Phi=15$ mila.

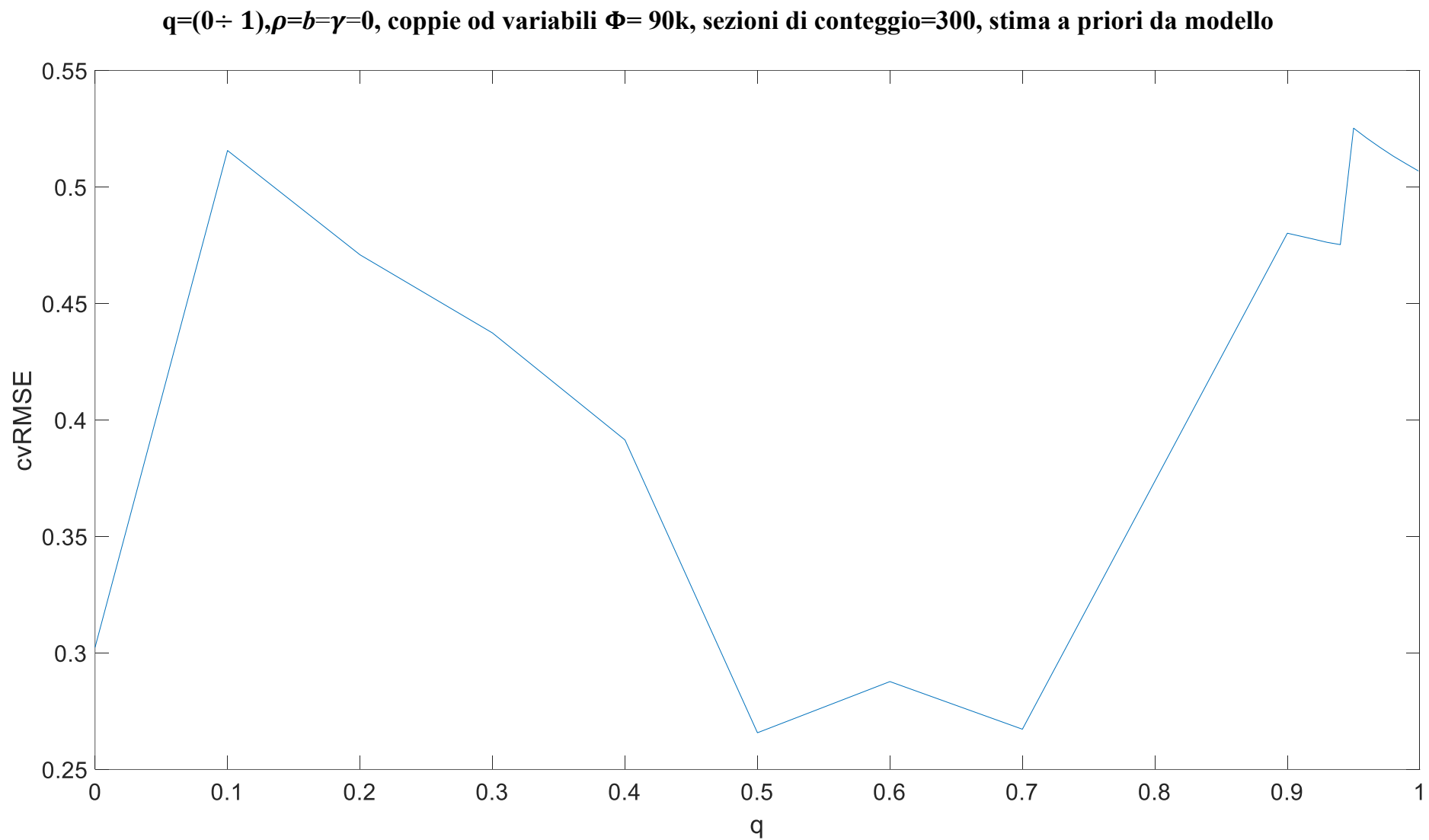


Fig. 5.19 cvRMSE hold out flussi d'arco, $q=(0 \div 1), \rho=b=\gamma=0$, coppie od variabili $\Phi=90k$, sezioni di conteggio=300, stima a priori da modello

La figura 5.19 rappresenta l'andamento dell'indicatore prestazionale $cvRMSE$, ottenuto attraverso l'ottimizzazione delle variabili x e del parametro $q \in [0,1]$, avendo preventivamente fissato $\rho=b=\gamma=0$, con tutte le coppie *od* variabili ($\Phi=90$ mila) e 300 sezioni di conteggio. Come si evince dalla figura, il modello di regressione che riesce a interpretare meglio la relazione di causa effetto tra la domanda e le variabili esplicative utilizzate è interposto tra un modello lineare e un modello moltiplicativo. Infatti, il valore del parametro q a cui compete il minimo errore rispetto ai flussi d'arco è pari a 0.5, in corrispondenza del quale il $cvRMSE$ assume il valore 0.26. I corrispondenti valore di $cvRMSE$ in condizione iniziale e a valle dell'applicazione della procedura di correzione classica sono rispettivamente 1.03 e 0.88. Dunque, rispetto alla procedura di correzione classica, la metodologia proposta contribuisce alla riduzione del 70% dell'errore.

6. Conclusioni e sviluppi futuri

L'elaborato di tesi ha affrontato il problema della correzione della stima dei flussi *od* mediante misure aggregate di traffico, con particolare riferimento al problema dello sbilancio tra equazioni (misure di traffico) e incognite (flussi *od* da stimare/correggere), perseguendo la strada della riduzione dimensionale delle incognite.

Il problema affrontato è uno dei più rilevanti nell'ambito della simulazione dei sistemi di trasporto, in quanto, i modelli di simulazione del traffico sono fortemente sensibili alle stime della domanda di mobilità (matrice *od*). Per questo motivo, tutti i principali software commerciali di simulazione del traffico disponibili sul mercato dispongono di routine per la correzione di matrici *od* mediante conteggi di traffico. Al contempo, pochi sono gli studi di letteratura che hanno evidenziato i limiti delle metodologie standard di correzione delle stime dei flussi *od* (es. Marzano, Papola and Simonelli, 2009, Simonelli et al. 2012). In sintesi, si può asserire che l'affidabilità delle procedure di correzione della stima della domanda che sfruttano misure aggregate di traffico sia legata alla posizione delle sezioni di conteggio e al rapporto tra equazioni linearmente indipendenti (conteggi) e incognite (flussi *od*) del problema. Nei casi reali, tale rapporto risulta essere molto basso e, di conseguenza, la correzione della stima della domanda potrebbe risultare scarsamente affidabile.

La presente tesi di dottorato ha investigato alcune procedure alternative a quelle di letteratura. Per perseguire tale obiettivo, sono state esplorate tecniche basate sull'aggregazione di incognite (capitolo 3 e capitolo 4) e tecniche fondate sulla costruzione di un modello dipendente da un numero ridotto di parametri e dalla introduzione di autocorrelazioni spaziali tra i flussi *od* (capitolo 5).

Nel capitolo 1 è vi è un'introduzione generale al problema della stima della domanda di trasporto, con una descrizione delle diverse metodologie di letteratura per la stima della

domanda, quali stima diretta, stima da modelli e correzione della domanda sulla base di conteggi di traffico.

Nel capitolo 2 si è presentato sinteticamente il problema della stima della domanda, ivi compreso quello della stima tramite conteggi di traffico, presentando la relativa letteratura scientifica rilevante. In particolare, sono state descritte nel dettaglio le tecniche di stima diretta della domanda, come le indagini campionarie a domicilio o a bordo veicolo, oppure le tecniche basate sull'utilizzo di traiettorie. Inoltre, sono state esposte le tecniche basate sull'utilizzo di modelli, che hanno l'obiettivo di riprodurre la domanda di trasporto in funzione di variabili esplicative e parametri strutturali, in virtù di una determinata forma funzionale (es. modelli derivati dal paradigma dei modelli di scelta discreta, modelli descrittivi di tipo gravitazionale etc). Infine, sono state descritte le tecniche di correzione della domanda di trasporto mediante l'utilizzo di dati aggregati di traffico, quali, ad esempio, i conteggi di traffico in sezioni predefinite della rete. Nello stesso capitolo, inoltre, sono stati anche descritti i principali problemi che rendono poco efficace la procedura di correzione della domanda di trasporto con conteggi di traffico, con riferimento particolare ai problemi relativi alla scarsa coverage e allo sbilancio equazioni-incognite. Per quanto riguarda il problema della scarsa coverage, è stato riportato lo stato dell'arte relativo alle metodologie attualmente presenti in letteratura per l'ottimizzazione del posizionamento delle sezioni di conteggio (*network sensor location*). Un'analoga analisi è stata condotta anche per il problema dello sbilancio equazioni-incognite. Infine, sono stati descritti i diversi stimatori che è possibile utilizzare per eseguire la correzione della domanda: minimi quadrati generalizzati, massima verosimiglianza e stimatori bayesiani.

Nel capitolo 3 è stata riportata la prima metodologia originale di riduzione dimensionale delle incognite del problema della correzione della domanda attraverso clustering sequenziale di flussi *od*, la quale si fonda su un'euristica sequenziale di aggregazioni di coppie *od*, nella quale viene ridotto di un'unità alla volta il numero di incognite del problema. L'applicazione della metodologia è stata testata su una toy network e sulla rete di Caserta, evidenziando ottimi risultati in termini di indicatori di errore rispetto ai valori veri della domanda e dei flussi di arco. I contenuti del presente capitolo sono rinvenibili in Buonocore *et al.* (2021).

Nel capitolo 4 è stata descritta una ulteriore metodologia finalizzata alla correzione della domanda di trasporto attraverso la riduzione dimensionale delle incognite, mediante clustering non sequenziale dei flussi *od*. In particolare, le tecniche di clustering adoperate sono state costruite ad hoc per la risoluzione del problema. Il numero totale di cluster è stato assunto coincidente al numero di sezioni di conteggio (numero di equazioni), in modo da riprodurre la condizione di *piena osservabilità*, dalla quale deriva l'unicità della soluzione al problema. La metodologia proposta è stata applicata alla rete della città di Torino, evidenziando ottimi risultati in termini di indicatori di errore. L'altro notevole vantaggio rispetto alla tecnica del clustering sequenziale risiede nella cospicua riduzione del tempo di elaborazione, che ha permesso l'applicazione della procedura su una rete significativamente più grande (90mila coppie *od* per la città di Torino contro le 4200 assunte per la città di Caserta) di quella analizzata nel cap. 3.

Nel capitolo 5, invece, è stata proposta una metodologia di correzione della domanda principalmente basata su una struttura del vettore di domanda dipendente da un numero di incognite ridotto. La metodologia impiega alcuni operatori della q -algebra (Borges, 2004; Nivanen et al., 2003), una forma di algebra deformata elaborata in supporto della statistica non estensiva di Tsallis (Tsallis, 1988). In particolare, l'utilizzo dell'operatore q -prodotto ha consentito di ottenere una forma funzionale di tipo gravitazionale generalizzato per i flussi di domanda, in funzione di un numero ridotto di variabili esplicative. Tale operatore generalizza gli operatori somma e prodotto in funzione di un parametro q da stimare sulla base di osservazioni (conteggi di traffico). L'ottimizzazione del parametro q consente di dedurre la relazione ottimale tra i flussi *od* e le variabili nello spazio ridotto, evidenziando quanto la relazione ottimale sia più vicina ad un modello di tipo additivo ($q=0$) o moltiplicativo ($q \rightarrow 1$). I risultati hanno mostrato che il valore ottimale di q risulta prossimo a 0.5, identificando così un legame che è perfettamente intermedio tra un modello additivo e un modello moltiplicativo. Da un punto di vista meramente analitico, il modello consente di generalizzare due procedure ben note nella letteratura relativa al tema di riduzione dimensionale, quali il modello gravitazionale semplice per la stima dei flussi di domanda da conteggi di traffico (Robillard, 1975) e una procedura analoga all'analisi fattoriale per la riduzione delle incognite (Djukic, Flötteröd, et al., 2012). Le tre metodologie proposte risultano utili a identificare possibili strade alternative alla procedura di correzione classica,

particolarmente fallace in tutte le situazioni reali in cui le condizioni risultano essere lontane da quelle ottimali di funzionamento (sbilancio equazioni-incognite elevato e ridotta coverage). Tutte e tre le metodologie proposte hanno fornito risultati molto promettenti e molto migliori rispetto alla procedura di correzione scelta come benchmark, cioè la procedura di correzione GLS proposta da Cascetta (1984). In particolare, l'ultima metodologia proposta ha fornito risultati estremamente positivi, con una riduzione dell'errore sui flussi assegnati estremamente alto, un valore di tale errore basso in senso assoluto (intorno al 20%) e quasi insensibile rispetto all'errore iniziale introdotto dalla stima a priori.

Relativamente alla scelta dello stimatore GLS per la correzione della domanda, è necessario sottolineare che non è stata fatta a caso. Infatti, oltre ad essere largamente utilizzato in letteratura scientifica ha un notevole vantaggio rispetto agli altri stimatori (bayesiani, massima verosimiglianza) cioè, non necessita di alcuna ipotesi sulla distribuzione delle stime a priori della domanda e dei flussi d'arco. Tra l'altro c'è da dire che le altre tipologie di stimatori, sotto opportune ipotesi di distribuzione dei parametri (stime a priori, conteggi etc.), sono analoghi allo stimatore GLS.

Sviluppi futuri consistono innanzitutto nella generalizzazione dei risultati ottenuti, ampliando e completando il piano sperimentale eseguito, nonché sperimentando le metodologie proposte su dei casi reali. Inoltre, altri sviluppi futuri sul tema sono sicuramente rappresentati da ulteriori estensioni delle metodologie sin qui testate. In particolare, la procedura di clustering simultanea testata al cap. 4 (già estensione di quella testata nel cap. 3) può essere migliorata con l'ausilio di set più ricchi di variabili esplicative relative alle zone della rete e algoritmi di clustering più efficienti rispetto a quelli adoperati per ottenere il partizionamento simultaneo dei flussi *od*. Inoltre, l'analisi del modello gravitazionale generalizzato testato nel cap. 5 può essere completata mediante uno studio più esteso dei parametri ρ e q (valori negativi o maggiori di 1), lo studio di algoritmi ad hoc di ottimizzazione per la risoluzione del problema dei minimi quadrati generalizzati e l'utilizzo di indicatori alternativi di autocorrelazione spaziale. Le procedure di tecniche di riduzione dimensionale adoperate nell'ambito dell'apprendimento automatico (machine learning), con particolare riferimento a quelle adoperate nel campo del riconoscimento delle immagini, possono sicuramente costituire un futuro banco di prova per ulteriori contributi

sul tema della riduzione delle incognite nel problema di stima dei flussi *od* da conteggi di traffico. Infine, il lavoro sin qui condotto sotto l'ipotesi di regime stazionario può certamente essere esteso per considerare una ulteriore evoluzione della domanda di tipo intra-periodale.

7. Bibliografia

- Balakrishna, R., Koutsopoulos, H. N., & Ben-Akiva, M. (2005). Calibration and validation of dynamic traffic assignment systems. *16th International Symposium on Transportation and Traffic Theory*, 407–426.
- Barceló, J., Montero, L., Bullejos, M., Serch, O., & Carmona, C. (2013). A Kalman Filter Approach for Exploiting Bluetooth Traffic Data When Estimating Time-Dependent OD Matrices. *Journal of Intelligent Transportation Systems*, 17(2), 123–141. <https://doi.org/10.1080/15472450.2013.764793>
- Bianco, L., Cerrone, C., Cerulli, R., & Gentili, M. (2014). Locating sensors to observe network arc flows: Exact and heuristic approaches. *Computers & Operations Research*, 46, 12–22. <https://doi.org/10.1016/j.cor.2013.12.013>
- Bianco, L., Confessore, G., & Reverberi, P. (2001). A network based model for traffic sensor location with implications on O/D matrix estimates. *Transportation Science*, 35(1), 50–60. Scopus. <https://doi.org/10.1287/trsc.35.1.50.10140>
- Borges, E. P. (2004). A possible deformed algebra and calculus inspired in nonextensive thermostatics. *Physica A: Statistical Mechanics and Its Applications*, 340(1–3), 95–101. <https://doi.org/10.1016/j.physa.2004.03.082>
- Box, G. E. P., & Cox, D. R. (1964). An Analysis of Transformations. *Journal of the Royal Statistical Society. Series B (Methodological)*, 26(2), 211–252.
- Buonocore, C., Marzano, V., Tinessa, F., Simonelli, F., & Papola, A. (2021). Improving o-d Flows Updating through Aggregation of o-d Pairs: Methodological Formulation and Performance Analysis of a Heuristic Algorithm. *2021 7th International Conference on Models and Technologies for Intelligent Transportation Systems (MT-ITS)*, 1–6. <https://doi.org/10.1109/MT-ITS49943.2021.9529269>
- Caggiani, L., Dell’Orco, M., Marinelli, M., & Ottomanelli, M. (2012). A Metaheuristic Dynamic Traffic Assignment Model for O-D Matrix Estimation using Aggregate

- Data. *Procedia - Social and Behavioral Sciences*, 54, 685–695.
<https://doi.org/10.1016/j.sbspro.2012.09.786>
- Caggiani, L., Ottomanelli, M., & Sassanelli, D. (2013). A fixed point approach to origin–destination matrices estimation using uncertain data and fuzzy programming on congested networks. *Transportation Research Part C: Emerging Technologies*, 28, 130–141. <https://doi.org/10.1016/j.trc.2010.12.005>
- Cantelmo, G., Viti, F., Cipriani, E., & Marialisa, N. (2015). A Two-Steps Dynamic Demand Estimation Approach Sequentially Adjusting Generations and Distributions. *2015 IEEE 18th International Conference on Intelligent Transportation Systems*, 1477–1482. <https://doi.org/10.1109/ITSC.2015.241>
- Capar, I., Kuby, M., Leon, V. J., & Tsai, Y.-J. (2013). An arc cover–path-cover formulation and strategic analysis of alternative-fuel station locations. *European Journal of Operational Research*, 227(1), 142–151. <https://doi.org/10.1016/j.ejor.2012.11.033>
- Cascetta, E. (1984). Estimation of trip matrices from traffic counts and survey data: A generalized least squares estimator. *Transportation Research Part B*. [https://doi.org/10.1016/0191-2615\(84\)90012-2](https://doi.org/10.1016/0191-2615(84)90012-2)
- Cascetta, E. (2006). *Modelli per i sistemi di trasporto* (1st ed.). Utet.
- Cascetta, E., & Nguyen, S. (1988). A unified framework for estimating or updating origin/destination matrices from traffic counts. *Transportation Research Part B: Methodological*, 22(6), 437–455. [https://doi.org/10.1016/0191-2615\(88\)90024-0](https://doi.org/10.1016/0191-2615(88)90024-0)
- Cascetta, E., Papola, A., Marzano, V., Simonelli, F., & Vitiello, I. (2013). *Quasi-dynamic estimation of o–d flows from traffic counts: Formulation, statistical validation and performance analysis on real data*. <https://doi.org/10.1016/j.trb.2013.06.007>
- Castillo, E., Conejo, A. J., Menéndez, J. M., & Jiménez, P. (2008). The Observability Problem in Traffic Network Models. *Computer-Aided Civil and Infrastructure Engineering*, 23(3), 208–222. <https://doi.org/10.1111/j.1467-8667.2008.00531.x>
- Chen, A., Pravinvongvuth, S., Chootinan, P., Lee, M., & Recker, W. (2007). Strategies for Selecting Additional Traffic Counts for Improving O-D Trip Table Estimation. *Transportmetrica*, 3(3), 191–211. <https://doi.org/10.1080/18128600708685673>
- Daganzo, C. F., & Sheffi, Y. (1977). On stochastic models of traffic assignment. *Transportation Science*. <https://doi.org/10.1287/trsc.11.3.253>

- Dantsuji, T., Hoang, N. H., Zheng, N., & Vu, H. L. (2022). A novel metamodel-based framework for large-scale dynamic origin–destination demand calibration. *Transportation Research Part C: Emerging Technologies*, 136, 103545. <https://doi.org/10.1016/j.trc.2021.103545>
- Demissie, M. G., de Almeida Correia, G. H., & Bento, C. (2013). Intelligent road traffic status detection system through cellular networks handover information: An exploratory study. *Transportation Research Part C: Emerging Technologies*, 32, 76–88. <https://doi.org/10.1016/j.trc.2013.03.010>
- Djukic, T., Flötteröd, G., van Lint, H., & Hoogendoorn, S. (2012). *Efficient real time OD matrix estimation based on Principal Component Analysis*. 115–121. <https://doi.org/10.1109/ITSC.2012.6338720>
- Djukic, T., Lint, J. W. C. V., & Hoogendoorn, S. P. (2012). Application of Principal Component Analysis to Predict Dynamic Origin–Destination Matrices: *Transportation Research Record*. <https://doi.org/10.3141/2283-09>
- Ehlert, A., Bell, M. G. H., & Grosso, S. (2006). The optimisation of traffic count locations in road networks. *Transportation Research Part B: Methodological*, 40(6), 460–479. <https://doi.org/10.1016/j.trb.2005.06.001>
- Fei, X., Mahmassani, H. S., & Murray-Tuite, P. (2013). Vehicular network sensor placement optimization under uncertainty. *Transportation Research Part C: Emerging Technologies*, 29, 14–31. <https://doi.org/10.1016/j.trc.2013.01.004>
- Fotheringham, A., Brunsdon, C., & Charlton, M. (2002). Geographically Weighted Regression: The Analysis of Spatially Varying Relationships. *John Wiley & Sons*, 13.
- Gan Liping, Yang Hai, & Wong Sze Chun. (2005). Traffic Counting Location and Error Bound in Origin-Destination Matrix Estimation Problems. *Journal of Transportation Engineering*, 131(7), 524–534. [https://doi.org/10.1061/\(ASCE\)0733-947X\(2005\)131:7\(524\)](https://doi.org/10.1061/(ASCE)0733-947X(2005)131:7(524))
- Ge, Q., & Fukuda, D. (2016). Updating origin–destination matrices with aggregated data of GPS traces. *Transportation Research Part C: Emerging Technologies*, 69, 291–312. <https://doi.org/10.1016/j.trc.2016.06.002>

- Gundlegard, D., & Karlsson, J. M. (2009). Route classification in travel time estimation based on cellular network signaling. *2009 12th International IEEE Conference on Intelligent Transportation Systems*, 1–6. <https://doi.org/10.1109/ITSC.2009.5309692>
- He, S. (2013). A graphical approach to identify sensor locations for link flow inference. *Transportation Research Part B: Methodological*, 51, 65–76. <https://doi.org/10.1016/j.trb.2013.02.006>
- Huang, H., Cheng, Y., & Weibel, R. (2019). Transport mode detection based on mobile phone network data: A systematic review. *Transportation Research Part C: Emerging Technologies*, 101, 297–312. <https://doi.org/10.1016/j.trc.2019.02.008>
- Iqbal, Md. S., Choudhury, C. F., Wang, P., & González, M. C. (2014). Development of origin–destination matrices using mobile phone call data. *Transportation Research Part C: Emerging Technologies*, 40, 63–74. <https://doi.org/10.1016/j.trc.2014.01.002>
- Larijani, A. N., Olteanu-Raimond, A.-M., Perret, J., Brédif, M., & Ziemlicki, C. (2015). Investigating the Mobile Phone Data to Estimate the Origin Destination Flow and Analysis; Case Study: Paris Region. *Transportation Research Procedia*, 6, 64–78. <https://doi.org/10.1016/j.trpro.2015.03.006>
- López-Ospina, H., Cortés, C. E., Pérez, J., Peña, R., Figueroa-García, J. C., & Urrutia-Mosquera, J. (2021). A maximum entropy optimization model for origin-destination trip matrix estimation with fuzzy entropic parameters. *Transportmetrica A: Transport Science*, 0(0), 1–38. <https://doi.org/10.1080/23249935.2021.1913257>
- Maher, M. J. (1983). Inferences on trip matrices from observations on link volumes: A Bayesian statistical approach. *Transportation Research Part B: Methodological*, 17(6), 435–447. <https://ideas.repec.org/a/eee/transb/v17y1983i6p435-447.html>
- Marzano, V., Papola, A., & Simonelli, F. (2009). Limits and perspectives of effective O–D matrix correction using traffic counts. *Transportation Research Part C: Emerging Technologies*, 17(2), 120–132. <https://doi.org/10.1016/j.trc.2008.09.001>
- Marzano, V., Papola, A., Simonelli, F., & Papageorgiou, M. (2018). A Kalman Filter for Quasi-Dynamic o-d Flow Estimation/Updating. *IEEE Transactions on Intelligent Transportation Systems*, 19(11), 3604–3612. <https://doi.org/10.1109/TITS.2018.2865610>

- Ng, M. (2013). Partial link flow observability in the presence of initial sensors: Solution without path enumeration. *Transportation Research Part E: Logistics and Transportation Review*, 51, 62–66. <https://doi.org/10.1016/j.tre.2012.12.002>
- Nivanen, L., Le Méhauté, A., & Wang, Q. A. (2003). Generalized algebra within a nonextensive statistics. *Reports on Mathematical Physics*, 52(3), 437–444. [https://doi.org/10.1016/s0034-4877\(03\)80040-x](https://doi.org/10.1016/s0034-4877(03)80040-x)
- Ortúzar, J., & Willumsen, L. (2011). *Modelling Transport* (4th ed.). Wiley.
- Park, J., Murphey, Y. L., McGee, R., Kristinsson, J. G., Kuang, M. L., & Phillips, A. M. (2014). Intelligent Trip Modeling for the Prediction of an Origin–Destination Traveling Speed Profile. *IEEE Transactions on Intelligent Transportation Systems*, 15(3), 1039–1053. <https://doi.org/10.1109/TITS.2013.2294934>
- Prakash, A. A., Seshadri, R., Antoniou, C., Pereira, F. C., & Ben-Akiva, M. (2018). Improving Scalability of Generic Online Calibration for Real-Time Dynamic Traffic Assignment Systems. *Transportation Research Record*, 2672(48), 79–92. <https://doi.org/10.1177/0361198118791360>
- Prakash, A. A., Seshadri, R., Antoniou, C., Pereira, F. C., & Ben-Akiva, M. E. (2017). Reducing the Dimension of Online Calibration in Dynamic Traffic Assignment Systems. *Transportation Research Record*, 2667(1), 96–107. <https://doi.org/10.3141/2667-10>
- Robillard, P. (1975). Estimating the O-D matrix from observed link volumes. *Transportation Research*, 9(2), 123–128. [https://doi.org/10.1016/0041-1647\(75\)90049-0](https://doi.org/10.1016/0041-1647(75)90049-0)
- Simonelli, F., Marzano, V., Papola, A., & Vitiello, I. (2012). A network sensor location procedure accounting for o–d matrix estimate variability. *Transportation Research Part B: Methodological*, 46(10), 1624–1638. <https://doi.org/10.1016/j.trb.2012.08.007>
- Simonelli, F., Tinessa, F., Buonocore, C., & Pagliara, F. (2020). Measuring the reliability of methods and algorithms for route choice set generation: Empirical evidence from a survey in the naples metropolitan area. *Open Transportation Journal*, 14, 50–66. doi:10.2174/1874447802014010050

- Simonelli, F., Tinessa, F., Marzano, V., Papola, A., & Romano, A. (2019). Laboratory experiments to assess the reliability of traffic assignment map. *2019 6th International Conference on Models and Technologies for Intelligent Transportation Systems (MT-ITS)*, 1–9. <https://doi.org/10.1109/MTITS.2019.8883390>
- Tinessa, F. (2021). Closed-form random utility models with mixture distributions of random utilities: Exploring finite mixtures of qGEV models. *Transportation Research Part B: Methodological*, 146, 262–288. <https://doi.org/10.1016/j.trb.2021.02.004>
- Tinessa, F., Simonelli, F., Marzano, V., & Buonocore, C. (2020). Evaluating the choice behaviour of high-speed rail passengers in Italy: A latent class structure with alternative kernel models to the Multinomial Logit. *2020 IEEE International Conference on Environment and Electrical Engineering and 2020 IEEE Industrial and Commercial Power Systems Europe (EEEIC / I CPS Europe)*, 1–6. <https://doi.org/10.1109/EEEIC/ICPSEurope49358.2020.9160556>
- Tsallis, C. (1988). Possible generalization of Boltzmann-Gibbs statistics. *Journal of Statistical Physics*, 52(1–2), 479–487. <https://doi.org/10.1007/BF01016429>
- Umarov, S., Tsallis, C., & Steinberg, S. (2008). On a q-Central Limit Theorem Consistent with Nonextensive Statistical Mechanics. *Milan Journal of Mathematics*, 76(1), 307–328. <https://doi.org/10.1007/s00032-008-0087-y>
- Wang, H., Calabrese, F., Di Lorenzo, G., & Ratti, C. (2010). *Transportation mode inference from anonymized and aggregated mobile phone call detail records*. 318–323. <https://doi.org/10.1109/ITSC.2010.5625188>
- Yang, H., Iida, Y., & Sasaki, T. (1991). An analysis of the reliability of an origin-destination trip matrix estimated from traffic counts. *Transportation Research Part B: Methodological*, 25(5), 351–363. [https://doi.org/10.1016/0191-2615\(91\)90028-H](https://doi.org/10.1016/0191-2615(91)90028-H)
- Yang, H., Yang, C., & Gan, L. (2006). Models and algorithms for the screen line-based traffic-counting location problems. *Computers & Operations Research*, 33(3), 836–858. <https://doi.org/10.1016/j.cor.2004.08.011>
- Yang, H., & Zhou, J. (1998). Optimal traffic counting locations for origin–destination matrix estimation. *Transportation Research Part B: Methodological*, 32(2), 109–126. [https://doi.org/10.1016/S0191-2615\(97\)00016-7](https://doi.org/10.1016/S0191-2615(97)00016-7)