

Università degli Studi di *Pavia*
Dipartimento di *Ingegneria dell'Informazione*
Corso di Laurea Magistrale di *Sensoristica e Strumentazione*
Biomedica in Bioingegneria

Ottimizzazione dei parametri di The Virtual Brain mediante algoritmi genetici

Tesi di Laurea Magistrale

Relatrice:

Prof.ssa *Claudia Gandini*

Candidato:

Damiano Anelli

Correlatrici:

Dott.ssa *Roberta Maria*

Lorenzi

Dott.ssa *Marta Gaviraghi*

Anno Accademico 2025/2026

Indice

Abstract	1
1 Introduzione	3
1.1 Risonanza Magnetica	3
1.1.1 Il segnale in Risonanza Magnetica	4
1.1.2 Immagini e contrasto	9
1.1.3 Risonanza Magnetica di diffusione	11
1.1.4 Risonanza Magnetica funzionale e segnale BOLD	14
1.2 The Virtual Brain	17
1.2.1 Modello dinamico locale di Wong-Wang	20
1.2.2 Parametri biofisici e ottimizzazione del modello	22
1.3 Ottimizzazione dei parametri del modello	23
1.3.1 Ottimizzazione attuale per grid search	23
1.3.2 Algoritmo NSGA-II	23
1.3.3 Selezione del punto ottimo	26
2 Obiettivo della tesi	29
3 Materiali e Metodi	31
3.1 Atassia	31
3.2 Coorte sperimentale	32
3.3 Acquisizioni di Risonanza Magnetica	34
3.4 Rete Attentiva Ventrle	34
3.5 Pre-processing dei dati MRI	35
3.5.1 Pre-processing dei dati di diffusione	36
3.5.2 Pre-processing dei dati resting-state fMRI	37
3.6 Connettività	39
3.6.1 Connettività strutturale	39

3.6.2	Connettività funzionale empirica statica	40
3.6.3	Connettività funzionale empirica dinamica	42
3.7	The Virtual Brain	43
3.7.1	Modello di Wong-Wang	44
3.7.2	Simulazione dell'attività della Rete Attentiva Ventrale .	45
3.8	Ottimizzazione multi-parametrica	46
3.8.1	Grid-search	46
3.8.2	NSGA-II e definizione degli obiettivi	47
3.8.3	Selezione del punto ottimo	50
4	Risultati	52
4.1	Confronto tra metodi di ottimizzazione	53
4.1.1	Confronto tra grid search e NSGA-II basato su PCC e KS	53
4.1.2	Confronto tra metodo basato su PCC e KS e metodo feature-based	55
4.2	Distribuzione dei parametri ottimizzati	57
4.2.1	Metodo basato su PCC e KS	58
4.2.2	Metodo feature-based	60
4.2.3	Fronte di Pareto	62
4.2.4	Confronto tra gruppi clinici	67
5	Discussione	73
6	Conclusioni	78

Abstract

Le alterazioni neurologiche possono riflettersi in variazioni della connettività strutturale e funzionale, modificando il bilanciamento tra attività eccitatoria e inibitoria, ovvero l'equilibrio tra popolazioni neuronali che amplificano il segnale e quelle che lo riducono. Questo equilibrio è fondamentale per la stabilità delle reti cerebrali e per il mantenimento delle sue funzioni cognitive e attentive. In questo contesto si inserisce The Virtual Brain, una piattaforma neuroinformatica che simula la dinamica cerebrale su larga scala a partire da dati individuali di risonanza magnetica. TVB utilizza modelli generativi di mesoscala, nei quali l'attività delle diverse regioni cerebrali viene descritta attraverso parametri biofisici che regolano le proprietà dinamiche locali della rete. Poiché questi parametri non sono direttamente misurabili in vivo, essi vengono stimati mediante procedure di model inversion, nelle quali il modello viene ottimizzato per riprodurre il più possibile i segnali empirici osservati.

L'obiettivo dello studio è ottimizzare i parametri biofisici del modello cerebrale implementato in TVB, il modello di Wong-Wang. Questo modello descrive la dinamica locale delle regioni cerebrali attraverso l'interazione tra popolazioni neuronali eccitatorie e inibitorie. Per l'ottimizzazione multi-parametrica, attualmente ottenuta con un metodo di grid-search per coppie di parametri, in questa tesi è stato utilizzato l'algoritmo genetico NSGA-II. Questo algoritmo lavora su una popolazione di possibili soluzioni, ciascuna corrispondente a una diversa combinazione di tutti i parametri del modello, e ad ogni generazione seleziona e combina le soluzioni migliori rispetto a due o più funzioni obiettivo, mantenendo quelle non dominate. In questo modo è possibile ottenere per ciascun soggetto un fronte di Pareto costituito da soluzioni che rappresentano diversi compromessi tra gli obiettivi considerati.

Sono state utilizzate due strategie di ottimizzazione. La prima è basata sulla minimizzazione di due funzioni obiettivo, relative al confronto tra la connettività funzionale statica simulata ed empirica e tra le distribuzioni di con-

nettività funzionale dinamica simulata ed empirica. La connettività funzionale statica è stata stimata come matrice di correlazione tra i segnali BOLD delle diverse regioni, mentre la connettività funzionale dinamica è stata calcolata suddividendo il segnale BOLD in finestre temporali scorrevoli.

La seconda strategia invece si concentra sul calcolo di feature in grado di descrivere proprietà specifiche del segnale BOLD e della rete funzionale. Queste feature sono state estratte, mediante il software Orange, con una procedura di selezione basata sull'Information Gain, con lo scopo di individuare le caratteristiche più informative nel distinguere i gruppi clinici. Le feature selezionate sono state il centroide spettrale del segnale BOLD, il valore medio della FCD ed infine i nodi core e la densità, misure derivanti dalla teoria dei grafi applicate alla rete funzionale. Il centroide spettrale descrive la distribuzione in frequenza del segnale BOLD, il valore medio della FCD sintetizza il livello medio di similarità tra le configurazioni funzionali dinamiche, i nodi core sono le regioni della rete funzionale che appartengono al nucleo più centrale e maggiormente connesso della rete, mentre la densità misura quanto la rete è connessa nel complesso.

L'analisi è stata condotta su una coorte di 27 soggetti composta da controlli sani e da soggetti caratterizzati da sindrome di Joubert e forme di atassia lentamente progressiva. Per ciascuna strategia di ottimizzazione è stato ottenuto un fronte di Pareto soggetto-specifico, dal quale è stato selezionato un punto ottimo mediante un criterio basato sulla distanza dal punto ideale. I risultati mostrano una variabilità dei parametri sia tra soggetti sia tra gruppi clinici. Inoltre, il confronto tra le due strategie evidenzia che funzioni obiettivo differenti possono guidare l'algoritmo verso combinazioni di parametri differenti.

Nel complesso, il lavoro ha portato allo sviluppo di un approccio di ottimizzazione multiparametrica per modelli cerebrali implementati in TVB. Rispetto alla grid search applicata a coppie di parametri, l'approccio proposto consente di esplorare simultaneamente lo spazio parametrico. Inoltre, può essere applicato sia a metriche basate sul confronto tra matrici di connettività, più vicine alle procedure standard in TVB, sia a feature rappresentative della rete funzionale e del segnale BOLD. In questo modo è possibile includere anche informazioni direttamente estratte dal segnale e proprietà dinamiche della connettività, non considerate in un approccio basato esclusivamente sul confronto tra matrici di connettività funzionale.

Capitolo 1

Introduzione

1.1 Risonanza Magnetica

La Risonanza Magnetica (Magnetic Resonance Imaging, MRI) è una tecnica di imaging non invasiva utilizzata sia in ambito clinico sia nella ricerca neuroscientifica. Questa tecnica consente di ottenere immagini tridimensionali del corpo umano con elevata risoluzione spaziale e con un ottimo contrasto tra i tessuti molli, ovvero tessuti caratterizzati da un elevato contenuto di acqua. Inoltre, a differenza di altre tecniche di imaging non è invasiva, poiché non utilizza radiazioni ionizzanti, ma si basa sull'impiego di campi magnetici statici e di impulsi a radiofrequenza. Per questo motivo è particolarmente adatta anche a studi che richiedono acquisizioni ripetute nel tempo, purché vengano rispettati gli opportuni vincoli tecnici e di sicurezza [1].

La MRI rappresenta uno strumento fondamentale nel contesto delle neuroscienze perché permette di acquisire informazioni sia anatomiche sia funzionali. Le immagini anatomiche forniscono una descrizione dettagliata della struttura cerebrale, mentre tecniche più avanzate, come la risonanza magnetica di diffusione (diffusion Magnetic Resonance Imaging, dMRI) e risonanza magnetica funzionale (functional Magnetic Resonance Imaging, fMRI), permettono rispettivamente di studiare il movimento microscopico delle molecole d'acqua all'interno dei tessuti e le variazioni dell'attività cerebrale nel tempo. Queste informazioni costituiscono la base per la costruzione di modelli computazionali del cervello, come quelli implementati in The Virtual Brain.

1.1.1 Il segnale in Risonanza Magnetica

Dal punto di vista fisico, la MRI si basa sui principi della Risonanza Magnetica Nucleare, che descrive il comportamento di specifici nuclei atomici quando vengono posti in un campo magnetico esterno e stimolati mediante impulsi a radiofrequenza. La descrizione della risonanza magnetica coinvolge il formalismo della meccanica quantistica; tuttavia, è possibile illustrarne i principi fondamentali anche nell'ambito della meccanica classica. Ad ogni nucleo atomico si può associare una quantità misurabile detto numero di spin (I), cioè una proprietà quantistica intrinseca che determina il numero di possibili livelli energetici del nucleo in presenza di un campo magnetico esterno. Il numero di livelli energetici è pari a $2I+1$ e i nuclei con numero pari sia di protoni che di neutroni presentano generalmente spin nullo ($I=0$) e di conseguenza non sono utilizzabili per la generazione del segnale di risonanza magnetica. Per la risonanza magnetica risultano importanti i nuclei con $I=1/2$, caratterizzati quindi da due soli livelli energetici. Tra questi, il protone dell'idrogeno (1H) rappresenta il nucleo di maggiore interesse, poichè è molto abbondante nei tessuti biologici a causa dell'elevata presenza di acqua. Sebbene lo spin non debba essere interpretato come una rotazione classica della particella su se stessa, per comprendere in modo intuitivo i principi della MRI è utile immaginare il protone come un piccolo dipolo magnetico. Il momento magnetico del protone, indicato con μ , è proporzionale al suo momento angolare intrinseco $\mathbf{P}$ attraverso il rapporto giromagnetico γ , che è una costante caratteristica di ogni specie nucleare:

$$\mu = \gamma \mathbf{P} \quad (1.1)$$

In assenza di un campo magnetico esterno, i momenti magnetici dei protoni sono orientati in modo casuale. Di conseguenza, la loro somma vettoriale è nulla e non si osserva una magnetizzazione macroscopica. Quando invece il campione viene posto all'interno di un forte campo magnetico statico, indicato con B_0 , i momenti magnetici tendono ad allinearsi parallelamente o antiparallelamente alla direzione del campo che, per convenzione, è orientato lungo l'asse z .

Nell'imaging medico per il campo B_0 si usano generalmente intensità da 1.5 T fino a 3 T ed è generato da un magnete superconduttivo, in modo da ottenere un campo il più possibile stabile e omogeneo. In presenza di B_0 , i protoni non

rimangono fermi, ma compiono un moto di precessione attorno alla direzione del campo magnetico. La frequenza angolare con la quale avviene questo moto è detta frequenza di Larmor ed è descritta dalla seguente equazione:

$$\omega_0 = \gamma B_0 \quad (1.2)$$

dove ω_0 è la frequenza angolare di precessione, γ è il rapporto giromagnetico e B_0 è l'intensità del campo magnetico statico.

A livello quantistico, l'applicazione del campo magnetico determina una separazione dei livelli energetici degli spin, fenomeno noto come effetto Zeeman. Per il protone, caratterizzato da numero quantico di spin $I = 1/2$, i possibili livelli energetici sono due, corrispondenti ai valori $m_I = [+1/2, -1/2]$, quindi uno a energia più bassa, associato all'orientamento parallelo al campo B_0 , e uno a energia più alta, associato all'orientamento antiparallelo. Poiché un numero leggermente maggiore di protoni occupa lo stato a energia più bassa, si genera una magnetizzazione netta lungo la direzione del campo magnetico statico, come illustrato in Figura 1.1.

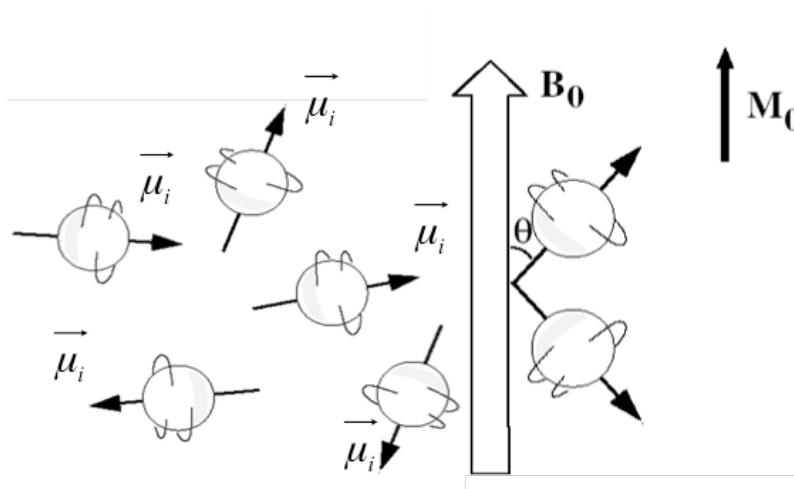


Figura 1.1: Orientamento dei momenti magnetici nucleari in assenza e presenza di un campo magnetico statico B_0 . In assenza di campo esterno, i momenti magnetici sono orientati casualmente e la magnetizzazione macroscopica risultante è nulla. In presenza di B_0 , i momenti magnetici tendono ad allinearsi lungo la direzione del campo, generando una magnetizzazione netta M_0 .

La magnetizzazione netta, indicata con $\mathbf{M}$, può essere definita come la somma vettoriale dei momenti magnetici dei singoli nuclei per unità di volume:

$$\mathbf{M} = \frac{1}{V} \sum_i \boldsymbol{\mu}_i \quad (1.3)$$

dove V è il volume considerato e $\boldsymbol{\mu}_i$ è il momento magnetico del singolo nucleo.

In condizioni di equilibrio la magnetizzazione è allineata con il campo B_0 ed è quindi costante nel tempo. Per questo motivo, da sola non genera un segnale misurabile. Per produrre un segnale di risonanza magnetica, è necessario perturbare il sistema applicando un impulso a radiofrequenza, indicato con $B_1(t)$, perpendicolare a B_0 . Quando la frequenza dell'impulso coincide con la frequenza di Larmor, si verifica la condizione di risonanza: i protoni assorbono energia e la magnetizzazione netta viene ruotata dalla direzione longitudinale verso il piano trasversale.

L'angolo di rotazione della magnetizzazione è chiamato flip angle e dipende dall'intensità e dalla durata dell'impulso a radiofrequenza. Un esempio comune è l'impulso a 90° , che porta la magnetizzazione nel piano trasversale. La componente trasversale della magnetizzazione ruota nel tempo alla frequenza di Larmor e induce una tensione nella bobina ricevente. Questo segnale elettrico rappresenta il segnale di risonanza magnetica misurabile.

Dopo l'applicazione dell'impulso a radiofrequenza, il sistema tende progressivamente a tornare alla condizione di equilibrio. Questo processo è chiamato rilassamento e descrive il recupero della magnetizzazione longitudinale e il decadimento della magnetizzazione trasversale. I due processi principali sono il rilassamento longitudinale, caratterizzato dalla costante di tempo T_1 , e il rilassamento trasversale, caratterizzato dalla costante di tempo T_2 [1].

Il rilassamento longitudinale, o rilassamento spin-reticolo, descrive il recupero della componente della magnetizzazione lungo l'asse z . Esso è legato allo scambio di energia tra gli spin e l'ambiente molecolare circostante. Dopo un impulso a 90° , la magnetizzazione longitudinale cresce progressivamente fino a tornare al valore di equilibrio M_0 , secondo un andamento esponenziale:

$$M_z(t) = M_0 (1 - e^{-t/T_1}) \quad (1.4)$$

La costante T_1 rappresenta il tempo necessario affinché la magnetizzazione longitudinale recuperi circa il 63% del suo valore di equilibrio.

Il rilassamento trasversale, o rilassamento spin-spin, descrive invece la per-

data progressiva della magnetizzazione nel piano trasversale. Questo fenomeno è dovuto alla perdita di coerenza di fase tra gli spin. Subito dopo l'impulso a radiofrequenza, gli spin precessano in fase; con il passare del tempo, però, interazioni locali e piccole differenze nel campo magnetico causano una progressiva dispersione delle fasi. La magnetizzazione trasversale decresce quindi secondo un andamento esponenziale:

$$M_{xy}(t) = M_{xy}(0)e^{-t/T_2} \quad (1.5)$$

La costante T_2 indica il tempo necessario affinché la magnetizzazione trasversale si riduca a circa il 37% del valore iniziale.

I tempi T_1 e T_2 dipendono dalle proprietà fisiche del tessuto. Per questo motivo, tessuti diversi mostrano tempi di rilassamento differenti. Questa caratteristica è alla base del contrasto nelle immagini di risonanza magnetica. Ad esempio, sostanza grigia, sostanza bianca e liquido cerebrospinale hanno valori diversi di T_1 e T_2 , e possono quindi essere distinti attraverso opportune sequenze di acquisizione.

Nella pratica il campo magnetico statico B_0 non è perfettamente omogeneo, poichè piccole imperfezioni del magnete e differenze di suscettibilità magnetica tra i tessuti generano variazioni locali del campo magnetico. Di conseguenza, gli spin in posizioni diverse possono precessare a frequenze leggermente differenti, contribuendo a una perdita di fase più rapida.

Il decadimento osservato della magnetizzazione trasversale non dipende quindi solo dal rilassamento T_2 , ma anche dalle disomogeneità del campo magnetico. L'effetto combinato viene descritto dalla costante T_2^* , che è sempre minore o uguale a T_2 . Possono essere uguali soltanto nel caso ideale in cui B_0 è perfettamente omogeneo e non ci sono disomogeneità locali. La relazione può essere espressa come:

$$\frac{1}{T_2^*} = \frac{1}{T_2} + \frac{1}{T_2'} \quad (1.6)$$

dove T_2' rappresenta il contributo dovuto alle disomogeneità del campo magnetico.

Dopo l'impulso a radiofrequenza, il segnale misurato prende il nome di free induction decay (FID). Il FID è generato dalla precessione della magnetizzazione trasversale e si manifesta come un segnale oscillante che decresce esponenzialmente nel tempo, a causa della progressiva perdita di magnetizza-

zione trasversale. Poiché nella pratica sono sempre presenti disomogeneità di campo, il FID decade secondo la costante T_2^* .

Per compensare almeno in parte la perdita di fase dovuta alle disomogeneità del campo magnetico, è possibile utilizzare un impulso di rifocalizzazione a 180° , applicato dopo l'impulso iniziale a 90° . Questa sequenza permette di generare un eco del segnale ed è nota come sequenza spin-echo, illustrata in figura 1.3.

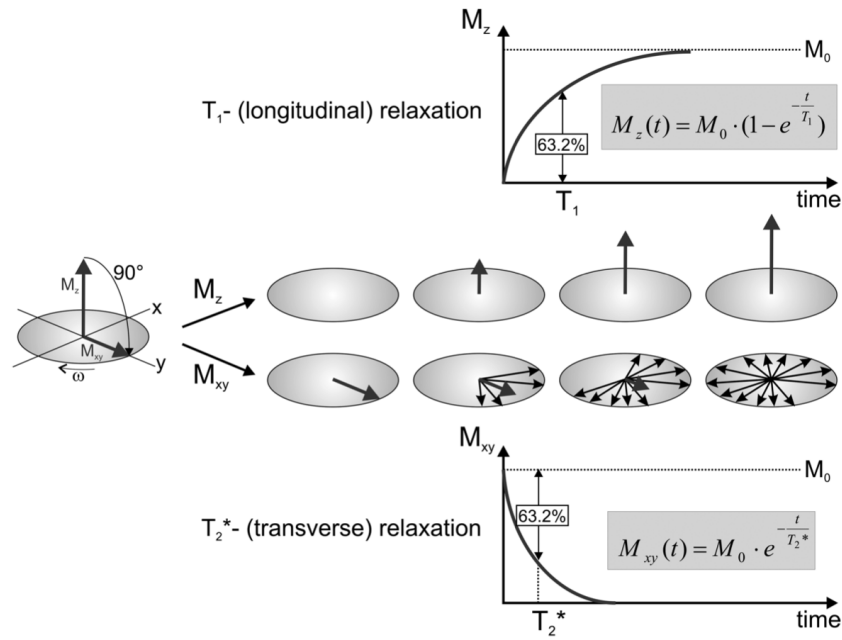


Figura 1.2: Rilassamento longitudinale T_1 e trasversale T_2^* dopo l'applicazione di un impulso a radiofrequenza. La componente longitudinale M_z tende progressivamente a recuperare il valore di equilibrio M_0 , mentre la componente trasversale M_{xy} decade nel tempo a causa della perdita di coerenza di fase degli spin. Jung e Weigel [2]

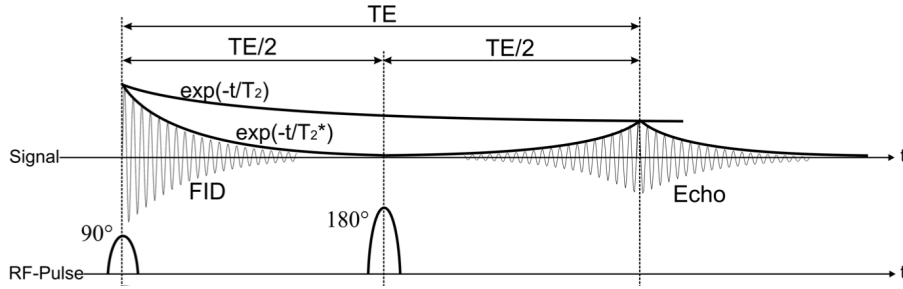


Figura 1.3: Sequenza spin-echo: dopo l'impulso a radiofrequenza di 90° , la magnetizzazione trasversale decade rapidamente generando il segnale FID, influenzato dal tempo di rilassamento T_2^* . L'applicazione di un impulso di rifocalizzazione a 180° al tempo $TE/2$ permette di compensare parzialmente la perdita di coerenza di fase degli spin e di generare un eco al tempo TE . Jung e Weigel [2]

1.1.2 Immagini e contrasto

Il segnale di risonanza magnetica raccolto dalla bobina ricevente proviene dall'intero volume eccitato. Per ricostruire un'immagine, è quindi necessario associare il segnale misurato alla sua posizione spaziale. Questo è possibile grazie all'utilizzo di gradienti di campo magnetico, cioè campi magnetici che variano linearmente nello spazio.

Quando viene applicato un gradiente lungo una direzione, i protoni localizzati in posizioni diverse sperimentano intensità di campo magnetico leggermente differenti. Di conseguenza, essi precessano a frequenze diverse, secondo la (1.2). La codifica spaziale viene ottenuta applicando gradienti di campo magnetico lungo le tre direzioni principali dello spazio. Il gradiente di selezione della fetta è un gradiente di campo magnetico che, applicato contemporaneamente a un impulso a radiofrequenza, consente di eccitare una specifica sezione del volume. Successivamente si applica il gradiente di codifica di fase. Questo gradiente di campo magnetico viene applicato per un breve intervallo di tempo prima dell'acquisizione del segnale e introduce una variazione di fase dipendente dalla posizione. Gli spin localizzati in punti diversi lungo la direzione di codifica di fase accumulano quindi fasi differenti, consentendo di distinguere la loro posizione spaziale. Infine si ha il gradiente di lettura, o codifica di frequenza, un gradiente di campo magnetico che viene applicato durante l'acquisizione del segnale e rende la frequenza di precessione degli spin dipendente dalla posizione lungo una determinata direzione spaziale: spin localizzati in

posizioni diverse precessano quindi a frequenze differenti. I dati acquisiti non rappresentano direttamente l'immagine finale, ma vengono inizialmente campionati nello spazio delle frequenze spaziali, chiamato k-spazio. Ogni punto del k-spazio contiene informazioni sulle frequenze spaziali dell'oggetto acquisito. Durante ogni acquisizione viene campionata una riga del k-spazio: il gradiente di lettura permette di acquisire i punti lungo tale riga, mentre il valore del gradiente di codifica di fase determina quale riga viene acquisita. Ripetendo la sequenza con diversi valori del gradiente di codifica di fase, il k-spazio viene progressivamente riempito. Le regioni centrali del k-spazio contengono principalmente informazioni sul contrasto globale dell'immagine e sono rappresentate da basse frequenze, mentre le regioni periferiche contengono informazioni sui dettagli e sui bordi e sono rappresentate da alte frequenze. L'immagine finale viene ricostruita applicando una trasformata di Fourier inversa ai dati acquisiti nel k-spazio.

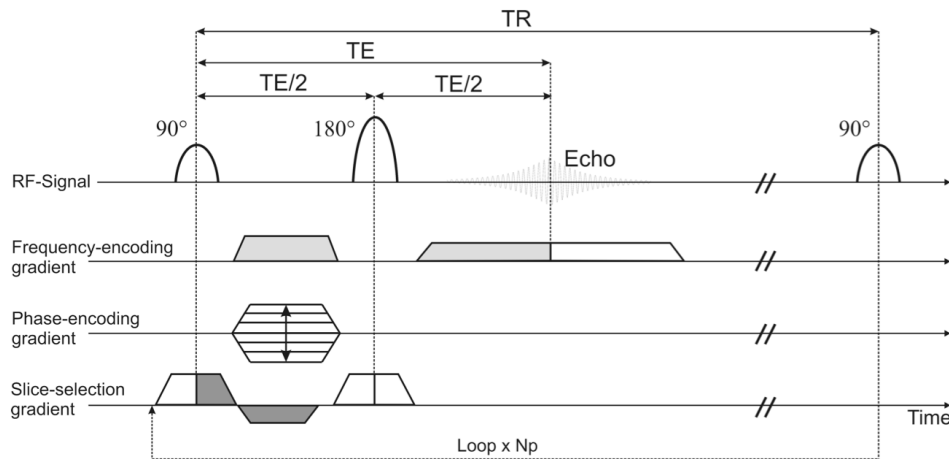


Figura 1.4: Schema temporale di una sequenza spin-echo per l'acquisizione MRI. L'impulso RF a 90° eccita gli spin, mentre l'impulso a 180° consente la formazione dell'eco al tempo TE. I gradienti di selezione della fetta, codifica di fase e codifica di frequenza permettono di associare il segnale acquisito alla posizione spaziale. La sequenza viene ripetuta più volte, variando il gradiente di codifica di fase, per riempire progressivamente le righe del k-spazio. Jung e Weigel [2].

In MRI, l'acquisizione del segnale è controllata da sequenze di eccitazione, ovvero schemi temporali che definiscono la successione di impulsi a radiofrequenza e gradienti di campo magnetico. La scelta della sequenza determina il

tipo di contrasto dell'immagine e il modo in cui vengono pesate le proprietà dei tessuti.

Due parametri fondamentali di una sequenza MRI sono il tempo di ripetizione (TR) e il tempo di echo (TE). Il TR è il tempo che intercorre tra due impulsi di eccitazione successivi, mentre il TE è il tempo tra l'impulso di eccitazione e la misura del segnale. Modificando TR e TE è possibile ottenere immagini pesate in densità protonica, T_1 o T_2 .

Le immagini pesate in densità protonica mettono in evidenza le differenze nel numero di protoni presenti nei tessuti. Le immagini pesate in T_1 enfatizzano le differenze nel recupero della magnetizzazione longitudinale, mentre le immagini pesate in T_2 evidenziano le differenze nel decadimento della magnetizzazione trasversale.

Tra le sequenze più utilizzate si trovano le sequenze spin-echo e gradient-echo. Le sequenze spin-echo utilizzano un impulso di rifocalizzazione a 180° per ridurre gli effetti delle disomogeneità di campo. Le sequenze gradient-echo, invece, non utilizzano l'impulso a 180° , ma generano l'eco attraverso una coppia di gradienti di verso opposto: il primo gradiente induce il defasamento degli spin, mentre il secondo gradiente ne produce il rifasamento parziale fino alla formazione dell'eco. Poiché questo meccanismo non compensa completamente le disomogeneità locali del campo magnetico, il segnale rimane sensibile al decadimento T_2^* .

1.1.3 Risonanza Magnetica di diffusione

La risonanza magnetica di diffusione (dMRI) è una tecnica che permette di ottenere informazioni sulla microstruttura dei tessuti analizzando il movimento delle molecole d'acqua. A differenza delle immagini anatomiche convenzionali, che descrivono principalmente differenze macroscopiche tra i tessuti, la dMRI sfrutta il fatto che l'acqua non si muove allo stesso modo in tutti gli ambienti biologici. Il movimento delle molecole d'acqua, infatti, dipende dalla presenza di membrane cellulari, fibre nervose, mielina e altre strutture microscopiche che possono ostacolare o indirizzare la diffusione.

In un mezzo libero e omogeneo, le molecole d'acqua tendono a muoversi casualmente in tutte le direzioni. Questo tipo di diffusione viene definito isotropo, poiché non esiste una direzione preferenziale di movimento. Nei tessuti biologici, invece, la diffusione può essere limitata dalla presenza di barriere

fisiche. In questo caso si parla di diffusione anisotropa, perché il movimento dell'acqua dipende dalla direzione considerata.

Dal punto di vista fisico, il processo diffusivo può essere descritto a partire dalla seconda legge di Fick, che mette in relazione la variazione temporale della concentrazione con la sua distribuzione spaziale. Nel caso monodimensionale, considerando un mezzo omogeneo e privo di barriere alla diffusione, la soluzione fondamentale assume la forma:

$$c(x, t) = \frac{1}{\sqrt{4\pi Dt}} e^{-\frac{x^2}{4Dt}} \quad (1.7)$$

dove $c(x, t)$ rappresenta la concentrazione delle molecole alla posizione (x) e al tempo (t) , mentre (D) indica il coefficiente di diffusione. Questa espressione descrive come, in un mezzo libero, la distribuzione delle molecole tenda ad allargarsi progressivamente nel tempo.

La possibilità di misurare la diffusione dell'acqua in MRI si basa sull'utilizzo di gradienti di diffusione, cioè gradienti di campo magnetico applicati per rendere il segnale sensibile al movimento microscopico delle molecole d'acqua. Nelle sequenze pesate in diffusione, comunemente indicate come diffusion weighted imaging (DWI), questi gradienti vengono applicati tipicamente in coppia. Il primo gradiente introduce uno sfasamento degli spin dipendente dalla loro posizione; successivamente viene applicato un secondo gradiente, che tende a rifasare gli spin.

Se gli spin rimangono nella stessa posizione tra l'applicazione dei due gradienti di diffusione, la fase accumulata durante il primo gradiente viene compensata dal secondo. Questa condizione corrisponde a una situazione di bassa diffusione in cui il segnale viene preservato. Al contrario, se le molecole d'acqua si spostano in quantità non trascurabile tra l'applicazione dei due gradienti, gli spin si trovano in una posizione diversa durante l'applicazione del secondo gradiente. In questo caso la rifocalizzazione della fase è incompleta: siamo in condizioni di elevata diffusione e si osserva una riduzione dell'ampiezza del segnale.

L'attenuazione del segnale è quindi legata alla quantità di diffusività del tessuto: maggiore è lo spostamento delle molecole d'acqua, maggiore è la perdita di coerenza di fase e minore è l'intensità del segnale misurato.

La sensibilità della sequenza alla diffusione è controllata dal cosiddetto fattore b , che dipende dall'intensità, dalla durata e dalla separazione temporale

dei gradienti di diffusione:

$$b = \gamma^2 G^2 \delta^2 \left(\Delta - \frac{\delta}{3} \right) \quad (1.8)$$

Valori di b più elevati rendono il segnale più sensibile al movimento delle molecole d'acqua, ma possono anche ridurre il rapporto segnale-rumore. Acquisendo immagini con almeno un'immagine non pesata in diffusione, indicata come $b = 0$, e una o più immagini pesate in diffusione, è possibile stimare il coefficiente di diffusione apparente (apparent diffusion coefficient, ADC). Questo parametro descrive la diffusione misurata all'interno del tessuto, che viene definita apparente perché non dipende solo dalla diffusione libera dell'acqua, ma anche dalle restrizioni imposte dalla microstruttura biologica.

Nel cervello la diffusione non può essere descritta completamente da un singolo valore scalare, poiché nella sostanza bianca il movimento dell'acqua è spesso direzionale. Per questo motivo si utilizza il tensore di diffusione (diffusion tensor imaging, DTI). Il DTI descrive la diffusione attraverso un tensore simmetrico, cioè una rappresentazione matematica tridimensionale che permette di stimare sia l'entità sia la direzione principale della diffusione in ciascun voxel. Dal tensore di diffusione è possibile ricavare diversi indici quantitativi, uno dei più rappresentativi è l'anisotropia frazionaria (fractional anisotropy, FA), che misura il grado di anisotropia della diffusione. Valori bassi di FA indicano una diffusione più isotropa, mentre valori elevati indicano una diffusione maggiormente orientata lungo una direzione preferenziale, come avviene nei principali fasci di sostanza bianca.

Le informazioni direzionali ottenute dalla dMRI possono essere utilizzate per ricostruire i fasci di sostanza bianca mediante trattografia, una tecnica di analisi delle immagini che permette di ricostruire virtualmente i fasci di fibre nervose. Questa tecnica non visualizza direttamente gli assoni, ma ricostruisce traiettorie, chiamate streamlines, seguendo localmente le direzioni principali di diffusione dell'acqua. In questo modo è possibile stimare le connessioni anatomiche tra diverse regioni cerebrali. La qualità della ricostruzione dipende da diversi fattori, tra cui la qualità dei dati di diffusione, il modello utilizzato per descrivere la diffusione, il tipo di metodo utilizzato per la trattografia e la parcellazione cerebrale adottata.

La dMRI assume un ruolo fondamentale perché permette di stimare la connettività strutturale del cervello. Dopo la parcellazione cerebrale, ogni re-

gione dell'atlante, ovvero una rappresentazione spaziale del cervello, può essere considerata come un nodo della rete, mentre le connessioni ricostruite tramite trattografia rappresentano i collegamenti anatomici tra i nodi. Da queste informazioni vengono costruite matrici di connettività strutturale, che possono includere il peso delle connessioni e la lunghezza media dei tratti. Queste matrici costituiscono l'input principale di The Virtual Brain, una piattaforma neuroinformatica che consente di simulare l'attività cerebrale su larga scala, poiché definiscono l'architettura anatomica lungo cui si propagano le dinamiche simulate dal modello.

In The Virtual Brain la dMRI non viene utilizzata solo come tecnica di imaging strutturale, ma come punto di partenza per costruire un modello computazionale personalizzato del cervello. La connettività strutturale derivata dalla dMRI vincola la comunicazione tra le regioni cerebrali simulate e influenza in modo diretto la dinamica funzionale prodotta dal modello. Per questo motivo, una ricostruzione corretta della rete strutturale è fondamentale per ottenere simulazioni confrontabili con i dati funzionali empirici.

1.1.4 Risonanza Magnetica funzionale e segnale BOLD

La risonanza magnetica funzionale (functional Magnetic Resonance Imaging, fMRI) è una tecnica di risonanza magnetica utilizzata per studiare in modo non invasivo le variazioni dell'attività cerebrale nel tempo. A differenza della MRI anatomica, che descrive principalmente la struttura dei tessuti, la fMRI fornisce informazioni funzionali misurando variazioni del segnale associate indirettamente all'attività neuronale.

La maggior parte degli studi fMRI si basa sul contrasto Blood Oxygenation Level Dependent (BOLD), che non richiede l'uso di mezzi di contrasto esterni, ma sfrutta le variazioni endogene dell'ossigenazione del sangue [3, 4, 5]. Dal punto di vista fisiologico, il segnale BOLD dipende dal legame tra attività neuronale, metabolismo e risposta vascolare. Quando una regione cerebrale aumenta la propria attività, cresce anche la richiesta locale di energia e ossigeno. In risposta, il sistema vascolare aumenta il flusso sanguigno cerebrale nella regione attiva. Questo aumento del flusso è generalmente superiore al consumo locale di ossigeno, producendo una variazione del rapporto tra emoglobina ossigenata e deossigenata nel sangue.

L'emoglobina può trovarsi principalmente in due forme: ossiemoglobina, quando è legata all'ossigeno, e deossiemoglobina, quando l'ossigeno è stato rilasciato ai tessuti. Queste due forme hanno proprietà magnetiche differenti. L'ossiemoglobina è diamagnetica e non altera il campo magnetico locale, mentre la deossiemoglobina è paramagnetica e introduce disomogeneità locali del campo magnetico. Tali disomogeneità accelerano il decadimento della magnetizzazione trasversale e riducono il tempo di rilassamento T_2^* , causando una diminuzione dell'intensità del segnale nelle immagini pesate per T_2^* [6].

Durante l'attivazione neuronale, l'aumento del flusso sanguigno porta a una riduzione relativa della concentrazione di deossiemoglobina nella regione attiva. Di conseguenza, le disomogeneità locali del campo magnetico diminuiscono e il segnale misurato nelle sequenze sensibili a T_2^* tende ad aumentare. Il segnale BOLD rappresenta quindi una misura indiretta dell'attività neuronale, mediata dalla risposta emodinamica e dall'accoppiamento neurovascolare, meccanismo attraverso cui l'attività neuronale induce variazioni locali del flusso sanguigno, ossigenazione e del volume ematico. Per questo motivo, la fMRI non misura direttamente l'attività elettrica dei neuroni, ma una risposta vascolare correlata all'attività cerebrale sottostante [7, 6].

La variazione temporale del segnale BOLD in risposta a un evento neuronale viene spesso descritta attraverso la hemodynamic response function (HRF), illustrata in figura 1.5.

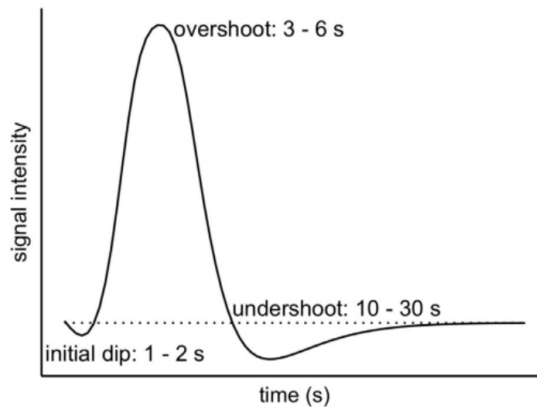


Figura 1.5: La risposta BOLD tipica, modellata attraverso la funzione di risposta emodinamica (Hemodynamic Response Function, HRF), mostra dopo l'inizio dello stimolo una breve riduzione iniziale del segnale, detta initial dip, associata al consumo locale di ossigeno. Successivamente il segnale aumenta, dando origine all'overshoot, che raggiunge il picco dopo circa 4–6 secondi. Infine, il segnale scende temporaneamente al di sotto del livello basale, fase nota come undershoot, prima di ritornare gradualmente alla baseline. Glover [8]

Dopo l'attivazione neuronale, il segnale non aumenta istantaneamente, ma segue una risposta più lenta legata alla dinamica vascolare. In una prima fase può essere osservata una lieve riduzione del segnale, detta initial dip, associata all'aumento iniziale del consumo di ossigeno e quindi a un incremento locale della deossiemoglobina. Successivamente, l'aumento del flusso sanguigno cerebrale supera la richiesta metabolica locale, determinando una riduzione relativa della deossiemoglobina, una diminuzione delle disomogeneità magnetiche locali e quindi un aumento del segnale BOLD fino al picco principale. Dopo il picco, il segnale ritorna gradualmente verso il livello di base e può presentare una fase di undershoot, durante la quale scende temporaneamente al di sotto del valore basale.

Dal punto di vista dell'acquisizione, la fMRI utilizza spesso sequenze Echo Planar Imaging (EPI), che permettono di acquisire rapidamente immagini dell'intero cervello con una risoluzione temporale compatibile con le variazioni lente del segnale BOLD. Poiché il contrasto BOLD dipende dal decadimento T_2^* , le acquisizioni fMRI sono generalmente basate su sequenze gradient-echo sensibili alle disomogeneità locali del campo magnetico.

Oltre agli studi basati su compiti specifici, la fMRI può essere acquisita anche in condizioni di riposo, senza che il soggetto esegua un task esplicito.

Questa modalità è nota come resting-state fMRI. In questo caso, l'interesse non è rivolto alla risposta a uno stimolo, ma alle fluttuazioni spontanee del segnale BOLD nel tempo. Le correlazioni temporali tra i segnali BOLD estratti da diverse regioni cerebrali permettono di stimare la connettività funzionale, cioè il grado di sincronizzazione statistica tra regioni cerebrali distinte [9].

La connettività funzionale viene comunemente rappresentata attraverso una matrice di correlazione, in cui ogni elemento descrive la relazione statistica tra due serie temporali BOLD regionali. A differenza della connettività strutturale, che descrive le connessioni anatomiche derivate dalla dMRI, la connettività funzionale riflette l'organizzazione dinamica dell'attività cerebrale. Due regioni possono mostrare un'elevata correlazione funzionale anche in assenza di una connessione anatomica diretta, poiché la loro attività può essere mediata da percorsi indiretti o da interazioni di rete più complesse.

Un'estensione della connettività funzionale descritta sopra è rappresentata dalla connettività funzionale dinamica (Functional Connectivity Dynamics, FCD), che descrive come i pattern di connettività funzionale cambiano nel tempo. Invece di calcolare una singola matrice di connettività sull'intera acquisizione, la FCD considera finestre temporali successive e confronta tra loro le configurazioni di connettività ottenute in diversi intervalli temporali. Questo permette di caratterizzare la variabilità temporale della dinamica cerebrale e il repertorio di stati funzionali esplorati dal cervello.

La fMRI ha un ruolo centrale perché fornisce il segnale BOLD empirico utilizzato per valutare la qualità delle simulazioni prodotte da The Virtual Brain. In particolare, dal BOLD empirico possono essere estratte misure di connettività funzionale, dinamica della connettività funzionale e altre feature descrittive della dinamica cerebrale. Queste grandezze vengono poi confrontate con le corrispondenti misure ottenute dalle simulazioni di dinamica cerebrale, contribuendo alla definizione degli obiettivi dell'ottimizzazione dei parametri del modello.

1.2 The Virtual Brain

The Virtual Brain (TVB) è una piattaforma neuroinformatica di simulazione dell'attività cerebrale soggetto-specifica su larga scala (e.g., segnali come il BOLD) a partire da dati individuali di risonanza magnetica. L'obiettivo di TVB è costruire un modello computazionale soggetto-specifico, in cui la strut-

tura anatomica del cervello viene utilizzata per simulare l'evoluzione temporale dell'attività neurale nelle diverse regioni cerebrali [10, 11].

In TVB il cervello viene rappresentato come una rete. I nodi della rete corrispondono alle regioni cerebrali definite da una parcellazione anatomica, mentre gli archi rappresentano le connessioni strutturali tra tali regioni. La connettività strutturale (structural connectivity, SC) viene ricostruita a partire da dati di risonanza magnetica di diffusione mediante trattografia ed è descritta da una matrice dei pesi, che rappresenta la forza relativa delle connessioni tra coppie di regioni, e da una matrice delle lunghezze, che descrive la lunghezza dei tratti che collegano i nodi della rete e definisce di conseguenza la distanza tra i nodi.

A ciascun nodo viene associato un modello dinamico locale, che descrive l'attività media della popolazione neuronale contenuta nella regione considerata. TVB non simula quindi l'attività dei singoli neuroni in un dominio di microscala, ma descrive l'attività cerebrale a un livello macroscopico, associando un modello di mesoscala a ogni nodo della rete considerata. In questo modo, l'attività di ciascuna regione dipende sia dalle proprietà del modello locale associato al nodo, sia dagli input provenienti dalle altre regioni attraverso la connettività strutturale. Il modello permette quindi di generare segnali cerebrali simulati, come il BOLD, che possono essere confrontati con i dati fMRI sperimentali per ottenere una descrizione soggetto-specifica della dinamica neuronale.

Gli input richiesti da TVB rappresentano le informazioni necessarie per costruire un modello personalizzato del cervello di ciascun soggetto. Un primo input è rappresentato dalle immagini anatomiche T1-pesate, utilizzate per definire lo spazio anatomico del soggetto e per applicare una parcellazione cerebrale. La parcellazione permette di suddividere la corteccia e le regioni sottocorticali in un insieme di aree distinte, definite secondo un atlante di riferimento. Queste regioni costituiscono i nodi della rete a cui verrà associato un modello di mesoscala. Un secondo input fondamentale è costituito dai dati di DWI, utilizzati per ricostruire la connettività strutturale del soggetto. Attraverso la trattografia è possibile stimare le connessioni anatomiche tra le diverse regioni cerebrali, ottenendo una matrice dei pesi, che descrive la forza relativa delle connessioni, e una matrice delle lunghezze o distanze, che rappresenta la lunghezza media dei tratti che collegano le varie regioni. Queste matrici definiscono l'architettura della rete cerebrale dove l'attività simulata

si propaga. I dati funzionali, quindi il segnale BOLD, generalmente acquisiti tramite resting-state fMRI, vengono invece utilizzati per ottenere un riferimento empirico della dinamica cerebrale del soggetto. A partire dalle serie temporali BOLD preprocessate è possibile calcolare la connettività funzionale empirica (indicata come *experimental functional connectivity*, *expFC*), che misura la correlazione di Pearson temporale tra l'attività BOLD delle diverse regioni cerebrali. Questi dati vengono utilizzati per valutare la bontà del modello, confrontando la connettività funzionale simulata (*simulated functional connectivity*, *simFC*), con quella osservata sperimentalmente.

Oltre ai dati di neuroimaging, TVB richiede la scelta di un modello dinamico locale e di un insieme di parametri biofisici iniziali. Il modello locale descrive l'evoluzione temporale dell'attività media di ciascun nodo, mentre i parametri biofisici regolano proprietà come la forza dell'accoppiamento globale, il bilanciamento tra eccitazione e inibizione e le dinamiche sinaptiche. Questi parametri vengono successivamente ottimizzati per rendere la dinamica simulata il più possibile simile a quella empirica.

In sintesi, i dati anatomici e di connettività strutturale definiscono la struttura del cervello simulato, mentre il modello dinamico locale e i parametri biofisici determinano l'evoluzione temporale dell'attività nei nodi della rete. I dati funzionali, invece, rappresentano il riferimento empirico con cui confrontare gli output della simulazione. Combinando queste informazioni, TVB permette di costruire un modello soggetto-specifico capace di simulare la comunicazione tra regioni cerebrali e di stimare parametri legati alla dinamica eccitatoria e inibitoria della rete.

L'output primario di TVB è costituito dalle serie temporali simulate dell'attività neurale in ciascun nodo. Poiché il dato sperimentale fMRI è rappresentato dal segnale BOLD, l'attività neurale simulata viene successivamente trasformata in un segnale BOLD simulato attraverso un modello emodinamico, in questo caso il modello Balloon-Windkessel, che va a modellare le variazioni di flusso sanguigno, volume ematico venoso e concentrazione di deossiemoglobina. In questo modo è possibile ottenere per ogni nodo una serie temporale BOLD simulata confrontabile con quella empirica. Da queste serie temporali si calcola la connettività funzionale simulata, ottenuta mediante il coefficiente di correlazione di Pearson tra le serie temporali BOLD simulate di ogni coppia di regioni cerebrali.

Il confronto tra *simFC* ed *expFC* rappresenta uno dei passaggi fonamen-

tali del fitting del modello. L'obiettivo della modellizzazione è avere una connettività funzionale generata dal modello che riproduce in modo più fedele la connettività funzionale osservata nei dati fMRI reali. Per questo motivo, i parametri del modello devono essere ottimizzati al fine di massimizzare la somiglianza tra dinamica simulata e dinamica empirica.

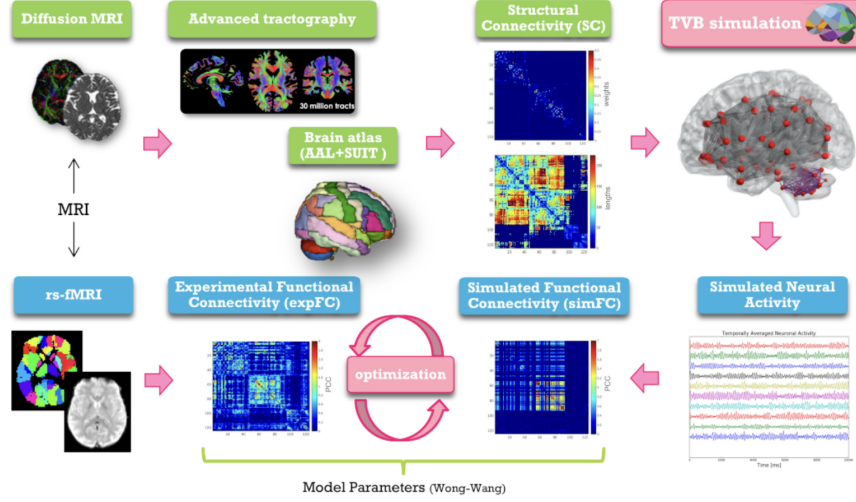


Figura 1.6: Procedimento di The Virtual Brain (TVB)

1.2.1 Modello dinamico locale di Wong-Wang

Uno dei modelli più utilizzati in TVB per simulare i dati di risonanza magnetica funzionale è il modello di Wong-Wang[12, 13, 14]. Questo modello descrive l'attività locale di ciascuna regione cerebrale come il risultato dell'interazione tra una popolazione neuronale eccitatoria e una popolazione neuronale inibitoria.

Nel modello, ogni nodo i della rete è rappresentato da due popolazioni: una eccitatoria, indicata con E , e una inibitoria, indicata con I . L'attività media di queste popolazioni è descritta attraverso variabili di gating sinaptico, $S_i^{(E)}$ e $S_i^{(I)}$, che rappresentano rispettivamente il livello medio di attivazione sinaptica eccitatoria e inibitoria nella regione i . Le popolazioni eccitatorie sono associate principalmente alla trasmissione mediata da recettori NMDA, mentre le popolazioni inibitorie sono associate alla trasmissione GABAergica.

Le correnti sinaptiche totali ricevute dalle popolazioni eccitatoria e inibitoria del nodo i sono descritte dalle seguenti equazioni:

$$I_i^{(E)} = W_E I_0 + w_p J_N S_i^{(E)} + G J_N \sum_j C_{ij} S_j^{(E)} - J_i S_i^{(I)} \quad (1.9)$$

$$I_i^{(I)} = W_I I_0 + J_N S_i^{(E)} - S_i^{(I)} \quad (1.10)$$

dove $I_i^{(E)}$ e $I_i^{(I)}$ rappresentano le correnti totali ricevute rispettivamente dalla popolazione eccitatoria e inibitoria del nodo i . Il termine I_0 indica l'input esterno costante, mentre W_E e W_I pesano tale input sulle due popolazioni. Il parametro J_N rappresenta la forza della trasmissione eccitatoria mediata dai recettori NMDA, J_i descrive la forza dell'inibizione locale, w_p regola la ricorrenza eccitatoria locale e G modula l'accoppiamento globale tra regioni cerebrali attraverso la matrice di connettività strutturale C_{ij} .

Le correnti sinaptiche vengono trasformate in firing rate mediante funzioni di trasferimento non lineari. Per la popolazione eccitatoria e inibitoria si ha:

$$r_i^{(E)} = H^{(E)} \left(I_i^{(E)} \right) = \frac{a_E I_i^{(E)} - b_E}{1 - \exp \left[-d_E \left(a_E I_i^{(E)} - b_E \right) \right]} \quad (1.11)$$

$$r_i^{(I)} = H^{(I)} \left(I_i^{(I)} \right) = \frac{a_I I_i^{(I)} - b_I}{1 - \exp \left[-d_I \left(a_I I_i^{(I)} - b_I \right) \right]} \quad (1.12)$$

dove $r_i^{(E)}$ e $r_i^{(I)}$ rappresentano i firing rate medi delle popolazioni eccitatoria e inibitoria nel nodo i , mentre a_E , b_E , d_E , a_I , b_I e d_I sono parametri che definiscono la forma delle funzioni di trasferimento.

L'evoluzione temporale delle variabili di gating sinaptico è descritta da un sistema di equazioni differenziali stocastiche:

$$\frac{dS_i^{(E)}(t)}{dt} = -\frac{S_i^{(E)}}{\tau_E} + \left(1 - S_i^{(E)} \right) \gamma_E r_i^{(E)} + \sigma v_i(t) \quad (1.13)$$

$$\frac{dS_i^{(I)}(t)}{dt} = -\frac{S_i^{(I)}}{\tau_I} + \gamma_I r_i^{(I)} + \sigma v_i(t) \quad (1.14)$$

dove τ_E e τ_I sono le costanti di tempo delle popolazioni eccitatoria e inibitoria, γ_E e γ_I sono fattori di guadagno sinaptico, σ rappresenta l'intensità del rumore e $v_i(t)$ è un termine stocastico. La presenza del rumore permette di simulare fluttuazioni spontanee dell'attività neurale, particolarmente importanti nel caso della resting-state fMRI.

Questa formulazione consente di descrivere la dinamica locale di ogni nodo come risultato del bilancio tra attività eccitatoria e inibitoria. In particolare, J_N e w_p contribuiscono alla componente eccitatoria del modello, mentre J_i

rappresenta il contributo inibitorio locale. Il parametro G , invece, controlla quanto la dinamica di ciascun nodo sia influenzata dalle connessioni strutturali con le altre regioni cerebrali.

1.2.2 Parametri biofisici e ottimizzazione del modello

Nel modello di Wong-Wang implementato in TVB, i parametri biofisici di interesse sono:

- G : accoppiamento globale, che regola la forza delle connessioni a lunga distanza tra regioni cerebrali;
- J_i : coupling inibitorio locale, associato alla forza dell'inibizione nel nodo;
- J_N : coupling eccitatorio mediato da recettori NMDA;
- w_p : ricorrenza eccitatoria locale, cioè la forza del feedback eccitatorio all'interno del nodo.

Il problema consiste quindi nell'individuare la combinazione di parametri che consente al modello di generare un'attività simulata il più possibile simile a quella osservata sperimentalmente.

Negli approcci classici, la bontà della simulazione viene spesso valutata confrontando la connettività funzionale simulata con quella empirica. Poiché l'obiettivo dell'ottimizzazione è minimizzare una funzione di costo, la correlazione può essere espressa come:

$$f_1 = 1 - PCC(\text{simFC}, \text{expFC}) \quad (1.15)$$

In questo modo, valori più bassi di f_1 corrispondono a una maggiore similarità tra la connettività funzionale simulata e quella empirica. Oltre alla connettività funzionale statica, è possibile considerare anche connettività funzionale dinamica (functional connectivity dynamics, FCD), che descrive come le relazioni funzionali tra regioni cambiano nel tempo. In questo caso, una misura utilizzata per confrontare la distribuzione dei valori di FCD simulati ed empirici è la statistica di Kolmogorov-Smirnov:

$$f_2 = KS(\text{FCD}_{sim}, \text{FCD}_{emp}) \quad (1.16)$$

Anche in questo caso, valori più bassi indicano una maggiore somiglianza tra simulazione e dati empirici.

Oltre all'approccio basato sulla minimizzazione delle due funzioni obiettivo $1 - PCC$ e KS , viene valutata una strategia alternativa di ottimizzazione basata su feature specifiche estratte dal segnale BOLD e dalle reti funzionali. Questa scelta è motivata dall'approccio suggerito dalla *Virtual Brain Inference* (VBI), che evidenzia l'utilità di considerare, oltre a misure di connettività funzionale come FC e FCD, anche caratteristiche legate al contenuto spettrale dei segnali simulati ed empirici. In aggiunta sono state incluse misure derivate dalla teoria dei grafi, poiché tali descrittori sono sensibili alle alterazioni della connettività nel caso specifico dell'atassia.

1.3 Ottimizzazione dei parametri del modello

In TVB la simulazione dell'attività cerebrale dipende dalla scelta dei parametri del modello. Poiché combinazioni diverse di parametri possono produrre segnali simulati più o meno coerenti con i dati sperimentali, è necessario introdurre una procedura di ottimizzazione. In questo lavoro, tale procedura è stata implementata mediante l'algoritmo genetico NSGA-II.

1.3.1 Ottimizzazione attuale per grid search

Prima dell'introduzione dell'algoritmo NSGA-II, l'ottimizzazione dei parametri del modello veniva effettuata mediante una procedura di *grid search*. Questo approccio consiste nel definire una griglia di valori per i parametri di interesse e nel valutare sistematicamente tutte le possibili combinazioni. Per ciascuna configurazione di parametri vengono eseguite le simulazioni in TVB e calcolate le funzioni obiettivo, utilizzate per confrontare i dati simulati con quelli empirici.

1.3.2 Algoritmo NSGA-II

NSGA-II, acronimo di *Non-dominated Sorting Genetic Algorithm II*, è un algoritmo evolutivo sviluppato per la risoluzione di problemi di ottimizzazione multi-obiettivo [15]. A differenza degli approcci a singolo obiettivo, nei quali si ricerca una sola soluzione ottimale, l'ottimizzazione multi-obiettivo mira a

individuare un insieme di soluzioni che rappresentano diversi compromessi tra funzioni obiettivo potenzialmente in conflitto. Questo insieme è definito fronte di Pareto.

Il concetto alla base di NSGA-II è la dominanza di Pareto. Una soluzione si dice dominante rispetto a un'altra se risulta migliore o uguale in tutti gli obiettivi considerati e strettamente migliore in almeno uno di essi. Le soluzioni che non sono dominate da nessun'altra soluzione della popolazione vengono considerate non dominate e costituiscono il primo fronte di Pareto. Le soluzioni dominate solo da quelle del primo fronte formano il secondo fronte, e così via. Questo procedimento prende il nome di *non-dominated sorting* e permette di assegnare a ogni individuo un rango in base alla qualità della soluzione.

NSGA-II lavora su una popolazione di soluzioni candidate che evolve nel corso di generazioni successive. A ogni generazione, gli individui vengono valutati rispetto alle funzioni obiettivo e ordinati in fronti di Pareto. Gli individui appartenenti ai fronti migliori hanno maggiore probabilità di essere selezionati per generare nuove soluzioni. In questo modo, l'algoritmo guida progressivamente la popolazione verso regioni dello spazio degli obiettivi caratterizzate da soluzioni migliori.

Un elemento fondamentale di NSGA-II è l'elitismo. A ogni iterazione, la popolazione dei genitori viene combinata con quella dei figli, formando una popolazione temporanea più grande. Da questa popolazione combinata vengono poi selezionati gli individui migliori per costruire la generazione successiva. Questo meccanismo permette di conservare le migliori soluzioni trovate nelle generazioni precedenti, evitando che vengano perse durante il processo evolutivo. La procedura di selezione elitista della nuova popolazione è rappresentata in Figura 1.7.

Oltre alla convergenza verso il fronte di Pareto, NSGA-II cerca di mantenere una buona diversità tra le soluzioni. Per questo utilizza la *crowding distance*, una misura che stima quanto una soluzione sia vicina alle altre nello spazio degli obiettivi. A parità di rango di Pareto, vengono preferite le soluzioni con *crowding distance* maggiore, cioè quelle localizzate in regioni meno affollate del fronte. Questo consente di ottenere un insieme di soluzioni più distribuito e rappresentativo dei diversi possibili compromessi.

La *crowding distance* viene calcolata separatamente per ciascun obiettivo ordinando le soluzioni in base al valore della funzione obiettivo considerata. Alle soluzioni estreme del fronte viene assegnata una distanza infinita, in modo

da preservare sempre i punti ai bordi del fronte di Pareto. Per le altre soluzioni, la distanza viene stimata come la somma delle distanze normalizzate tra le soluzioni vicine nello spazio degli obiettivi:

$$d_i = \sum_{m=1}^M \frac{f_{i+1}^{(m)} - f_{i-1}^{(m)}}{f_{\max}^{(m)} - f_{\min}^{(m)}} \quad (1.17)$$

dove M rappresenta il numero di obiettivi considerati, $f_{i+1}^{(m)}$ e $f_{i-1}^{(m)}$ sono i valori dell'obiettivo m -esimo delle soluzioni adiacenti alla soluzione i , mentre $f_{\max}^{(m)}$ e $f_{\min}^{(m)}$ rappresentano rispettivamente il valore massimo e minimo dell'obiettivo all'interno del fronte considerato.

Una crowding distance elevata indica che la soluzione si trova in una regione meno densa del fronte di Pareto e viene quindi favorita durante il processo di selezione. In questo modo, NSGA-II riesce a bilanciare simultaneamente convergenza verso il fronte ottimale e mantenimento della diversità delle soluzioni.

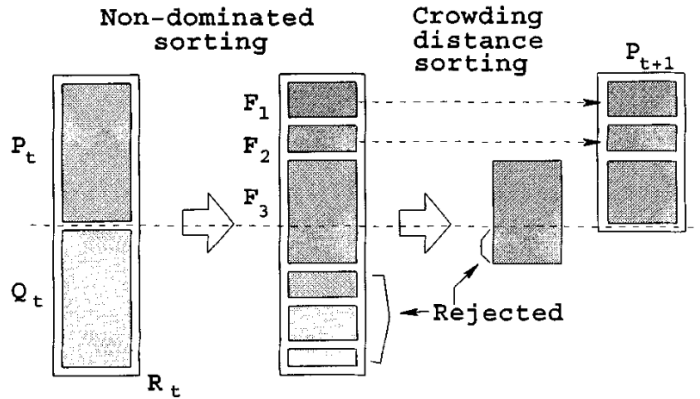


Figura 1.7: Procedura elitista di NSGA-II. Le popolazioni dei genitori P_t e dei figli Q_t vengono combinate nella popolazione R_t , ordinata in fronti non dominati. I fronti migliori vengono copiati in P_{t+1} ; se un fronte non può essere incluso interamente, vengono selezionate le soluzioni con maggiore *crowding distance*. Tratta da Deb et al. [15].

La generazione di nuove soluzioni avviene mediante operatori tipici degli algoritmi genetici, come selezione, crossover e mutazione. La selezione privilegia gli individui migliori secondo il rango di Pareto e la crowding distance, il crossover combina le caratteristiche di due soluzioni per crearne di nuove e la mutazione introduce piccole variazioni casuali, favorendo l'esplorazione dello spazio

di ricerca. L'equilibrio tra selezione delle soluzioni migliori ed esplorazione di nuove configurazioni permette all'algoritmo di migliorare progressivamente la qualità della popolazione.

1.3.3 Selezione del punto ottimo

Un'ottimizzazione multi-obiettivo non restituisce generalmente una sola soluzione finale, ma un insieme di soluzioni non dominate. Queste soluzioni costituiscono il fronte di Pareto e rappresentano configurazioni diverse dei parametri del modello, ciascuna associata a un diverso equilibrio tra gli obiettivi considerati. Per ogni soggetto e per ciascuna strategia di ottimizzazione, è stato quindi ottenuto un fronte di Pareto specifico, formato da diverse combinazioni dei parametri TVB.

Nelle analisi successive è stato necessario ricondurre ciascun fronte a una singola soluzione rappresentativa. Questo passaggio permette di confrontare i soggetti tra loro e di descrivere i risultati in modo più sintetico. La scelta di tale soluzione non è immediata, poiché in un problema multi-obiettivo non esiste un punto migliore in senso assoluto. Ogni soluzione del fronte rappresenta infatti un compromesso: il miglioramento di un obiettivo può essere associato al peggioramento di un altro. Per questo motivo, la selezione del punto rappresentativo deve essere effettuata mediante un criterio coerente con la struttura del fronte e con il numero di obiettivi utilizzati.

Criterio della distanza dal punto ideale

Questo approccio consiste nel selezionare la soluzione più vicina al punto ideale. In un problema di minimizzazione, il punto ideale corrisponde alla condizione teorica in cui tutti gli obiettivi assumono contemporaneamente il loro valore minimo. Tale condizione rappresenta quindi una soluzione ottimale dal punto di vista di tutte le funzioni obiettivo, ma nella pratica non è quasi mai presente all'interno del fronte di Pareto. Per questo motivo, si individua la soluzione reale che si avvicina maggiormente a tale riferimento.

Per applicare questo criterio, i valori delle funzioni obiettivo sono stati organizzati in una matrice F , in cui ogni riga corrisponde a una soluzione del fronte di Pareto e ogni colonna a uno specifico obiettivo. Poiché gli obiettivi possono essere espressi su scale numeriche diverse, è stata applicata una normalizzazione min-max a ciascuna colonna della matrice:

$$F_{ij}^{\text{norm}} = \frac{F_{ij} - \min(F_j)}{\max(F_j) - \min(F_j) + \epsilon}, \quad (1.18)$$

dove F_{ij} indica il valore dell'obiettivo j per la soluzione i , mentre ϵ è un termine numerico molto piccolo introdotto per evitare eventuali divisioni per zero. Nel codice utilizzato, tale termine è stato impostato pari a 10^{-12} .

Dopo la normalizzazione, tutti gli obiettivi risultano confrontabili sulla stessa scala. In questo spazio normalizzato, il punto ideale coincide con l'origine, poiché il valore migliore di ciascun obiettivo corrisponde a zero. Per ogni soluzione del fronte è stata quindi calcolata la distanza euclidea dall'origine:

$$d_i = \|F_i^{\text{norm}}\|_2 = \sqrt{\sum_{j=1}^m (F_{ij}^{\text{norm}})^2}, \quad (1.19)$$

dove m rappresenta il numero di obiettivi considerati. La soluzione selezionata è stata quindi quella caratterizzata dalla distanza minima:

$$i^* = \arg \min_i d_i. \quad (1.20)$$

Questo criterio consente di individuare il punto che si avvicina maggiormente alla minimizzazione simultanea degli obiettivi. Inoltre, l'utilizzo della normalizzazione riduce il rischio che la selezione sia influenzata dalla diversa scala numerica delle funzioni obiettivo.

Criterio del ginocchio

Un altro criterio possibile è quello del ginocchio, spesso indicato come *knee point*. Questo punto corrisponde a una regione del fronte in cui il compromesso tra gli obiettivi cambia in modo più marcato. In particolare, il knee point può essere interpretato come una soluzione in cui un ulteriore miglioramento di un obiettivo richiederebbe un peggioramento relativamente elevato di un altro. Per questo motivo, viene spesso considerato un buon candidato per rappresentare un compromesso bilanciato.

Nel caso di un fronte bidimensionale, il significato geometrico del knee point è abbastanza intuitivo. Si possono considerare i due estremi del fronte, cioè le soluzioni che ottimizzano rispettivamente il primo e il secondo obiettivo. La retta che congiunge questi due punti rappresenta una sorta di riferimento lineare tra le due soluzioni estreme. Il ginocchio del fronte può quindi essere

individuato come il punto che si trova alla massima distanza da tale retta. Questo punto identifica la porzione del fronte in cui la curvatura è più evidente e in cui il passaggio da un obiettivo all'altro risulta meno lineare.

Il principale vantaggio di questo criterio è la sua interpretabilità, soprattutto quando il fronte è formato da due soli obiettivi. Tuttavia, la sua applicazione diventa meno immediata quando il numero di obiettivi aumenta. Nel caso di fronti tridimensionali o multidimensionali, infatti, il concetto di ginocchio deve essere generalizzato, ad esempio considerando la distanza da un piano o da un iperpiano definito dagli estremi del fronte. Questa procedura può essere più sensibile alla distribuzione delle soluzioni, alla scelta dei punti estremi e alla presenza di eventuali valori anomali.

Capitolo 2

Obiettivo della tesi

Nei modelli computazionali cerebrali, una delle strategie più comunemente utilizzate per l'ottimizzazione parametrica è la grid-search. Questo approccio può però presentare alcuni limiti, soprattutto quando aumenta il numero di parametri da ottimizzare, poiché richiede un'elevata quantità di simulazioni e può rendere inefficiente l'esplorazione dello spazio dei parametri. Per questo motivo, si è deciso di testare un approccio alternativo basato su algoritmi genetici, più adatto per problemi multi-obiettivo.

L'obiettivo principale dello studio è confrontare i risultati ottenuti dall'ottimizzazione mediante NSGA-II con quelli derivanti da metodi basati su grid-search. Questo confronto consente di valutare se l'utilizzo di un algoritmo genetico permetta di ottenere valori migliori delle funzioni obiettivo rispetto alla strategia di esplorazione parametrica utilizzata in studi precedenti.

Un ulteriore obiettivo è ottenere simulazioni della dinamica BOLD il più possibile confrontabili con i dati empirici. L'ottimizzazione non viene considerata solo come una procedura numerica per migliorare il fitting del modello, ma anche come uno strumento per analizzare in modo quantitativo le proprietà dinamiche della rete cerebrale simulata.

Nello specifico si va a confrontare la connettività funzionale statica simulata ed empirica e le rispettive distribuzioni di connettività funzionale dinamica, mentre in un altro approccio si utilizzano feature estratte dal segnale BOLD e dalla rete funzionale, con l'obiettivo di guidare l'ottimizzazione verso proprietà specifiche della dinamica cerebrale.

Le feature considerate comprendono il centroide spettrale del segnale BOLD, il valore medio di FCD, i nodi core e la densità della connettività funzionale.

Queste misure sono state selezionate perché descrivono aspetti complementari della dinamica funzionale, includendo sia proprietà temporali del segnale sia caratteristiche topologiche della rete. In questo modo, l'approccio feature-based permette di valutare se informazioni più sintetiche, ma potenzialmente più descrittive della dinamica cerebrale, possano rappresentare un'alternativa efficace rispetto al confronto diretto tra matrici di connettività.

Infine, lo studio intende analizzare le distribuzioni parametriche ed esplorare se i parametri ottimizzati possano evidenziare differenze significative tra i gruppi clinici presenti nella coorte, costituita da controlli sani e da soggetti con alterazioni cerebellari. Questo confronto può contribuire a comprendere se TVB, combinata con tecniche di ottimizzazione multi-obiettivo, sia in grado di fornire informazioni utili sulla dinamica cerebrale e sulle possibili alterazioni dell'equilibrio eccitatorio/inibitorio nei diversi gruppi analizzati.

Capitolo 3

Materiali e Metodi

3.1 Atassia

Le atassie cerebellari comprendono un insieme eterogeneo di condizioni neurologiche accomunate da un'alterazione della struttura cerebellare. Dal punto di vista clinico, l'atassia si manifesta principalmente con difficoltà nella coordinazione del movimento, instabilità posturale e alterazioni dell'equilibrio, in assenza di un deficit primario di forza muscolare. Poiché il cervelletto partecipa non solo al controllo motorio, ma anche alla regolazione di funzioni cognitive, attentive e affettive, il danno cerebellare può inoltre associarsi a deficit neuropsicologici di diversa entità [16]. Sono stati considerati due gruppi patologici: soggetti con sindrome di Joubert e soggetti con atassia lentamente progressiva.

La sindrome di Joubert rappresenta una forma congenita e non progressiva di atassia cerebellare. Si tratta di una rara patologia genetica dello sviluppo neurologico, caratterizzata da alterazioni strutturali del cervelletto e del tronco encefalico. Dal punto di vista clinico, i pazienti con sindrome di Joubert possono presentare ritardo dello sviluppo motorio, atassia a esordio precoce e compromissione cognitiva. Trattandosi di una condizione congenita, il quadro tende generalmente a rimanere relativamente stabile nel tempo, con la possibilità di un parziale adattamento funzionale durante lo sviluppo [16].

Le atassie lentamente progressive (slowly progressive ataxia, SP) comprendono invece condizioni di origine genetica, metabolica o multisistemica caratterizzate da un deterioramento graduale della funzione cerebellare. A differenza delle forme congenite non progressive, le SP sono associate a una progressiva atrofia cerebellare e a un peggioramento clinico che può svilupparsi nel corso

di anni o decenni. Queste condizioni condividono manifestazioni cliniche comuni, tra cui alterazioni della coordinazione motoria, disturbi dell'equilibrio e possibile coinvolgimento cognitivo-affettivo [16].

Il confronto tra sindrome di Joubert e atassie lentamente progressive può risultare particolarmente rilevante in uno studio di modellistica cerebrale poichè nella sindrome di Joubert il danno cerebellare è presente fin dalle prime fasi dello sviluppo e può favorire meccanismi di riorganizzazione compensatoria, mentre nelle atassie progressive il sistema nervoso è sottoposto a un'alterazione continua e progressiva, che può limitare la possibilità di mantenere un equilibrio funzionale stabile. Per questo motivo, l'analisi delle reti cerebrali permette di indagare non solo la presenza di alterazioni locali del cervelletto, ma anche il modo in cui tali alterazioni si riflettono sulla comunicazione tra regioni cerebrali appartenenti a reti funzionali più ampi.

3.2 Coorte sperimentale

La coorte sperimentale è costituita da 27 soggetti, suddivisi in tre gruppi: 8 soggetti con sindrome di Joubert, 8 soggetti con atassia cerebellare lentamente progressiva e 11 controlli sani. I soggetti sono stati reclutati presso l'Istituto Neurologico Carlo Besta di Milano. Il gruppo patologico comprendeva soggetti con diagnosi geneticamente confermata di sindrome atassica cerebellare pediatrica e con anomalie cerebellari documentate mediante risonanza magnetica, quali ipoplasia o atrofia del verme, degli emisferi cerebellari o dell'intero cervelletto [16]. I controlli sani sono stati inclusi come gruppo di riferimento per il confronto con i due gruppi patologici. I criteri di esclusione comprendevano controindicazioni alla risonanza magnetica, presenza di altre patologie neurologiche o psichiatriche e, per i controlli, ulteriori fattori potenzialmente interferenti come fattori di rischio cardiovascolare o cerebrovascolare, storia di trauma cranico, uso di trattamenti psicoattivi o condizioni cliniche recenti potenzialmente rilevanti. Tutti i partecipanti sono stati inoltre invitati a evitare l'assunzione di farmaci nelle 24 ore precedenti l'acquisizione MRI.

La coorte è stata utilizzata per confrontare le proprietà dinamiche e di connettività funzionale della rete ventrale attentiva (VAN) nei tre gruppi clinici.

Le caratteristiche della coorte sono riportate in Tabella 3.1. Nel complesso, la coorte comprendeva 16 soggetti di sesso maschile e 11 soggetti di sesso femminile. L'età media complessiva era pari a circa 20.6 anni, con un intervallo

compreso tra 12 e 31 anni. Il gruppo HC comprendeva 6 maschi e 5 femmine, con età media pari a circa 20.5 anni. Il gruppo JS comprendeva 5 maschi e 3 femmine, con età media pari a circa 21.3 anni. Il gruppo SP comprendeva 5 maschi e 3 femmine, con età media pari a circa 20.1 anni. Queste informazioni sono importanti perché età e sesso possono contribuire alla variabilità inter-soggetto delle misure di connettività e dei parametri simulati, soprattutto in coorti numericamente limitate.

Tabella 3.1: Coorte sperimentale utilizzata per l'analisi della Ventral Attention Network.

Soggetto	Sesso	Età	Gruppo	Network
01	M	21	JS	Ventral attention
02	F	16	SP	Ventral attention
03	F	19	SP	Ventral attention
04	F	16	SP	Ventral attention
05	F	27	JS	Ventral attention
06	M	12	JS	Ventral attention
07	M	31	JS	Ventral attention
08	M	31	SP	Ventral attention
09	M	20	SP	Ventral attention
10	M	21	SP	Ventral attention
11	M	22	JS	Ventral attention
12	M	20	SP	Ventral attention
13	F	17	JS	Ventral attention
14	F	19	JS	Ventral attention
15	M	21	JS	Ventral attention
16	M	18	SP	Ventral attention
17	M	22	HC	Ventral attention
18	F	21	HC	Ventral attention
19	M	21	HC	Ventral attention
20	F	15	HC	Ventral attention
21	F	18	HC	Ventral attention
22	F	23	HC	Ventral attention
23	F	21	HC	Ventral attention
24	M	23	HC	Ventral attention
25	M	20	HC	Ventral attention
26	M	21	HC	Ventral attention
27	M	21	HC	Ventral attention

3.3 Acquisizioni di Risonanza Magnetica

Le immagini di risonanza magnetica sono state acquisite mediante uno scanner Philips Achieva a 3 T, utilizzando un protocollo standardizzato all'interno della rete italiana di neuroscienze e neuroriabilitazione. Il protocollo includeva tre acquisizioni principali: un'immagine anatomica T1-pesata tridimensionale, una sequenza di diffusione multi-shell e una sequenza funzionale resting-state [16]. L'immagine anatomica T1-pesata è stata acquisita mediante sequenza 3D Fast Field Echo. I principali parametri di acquisizione erano: tempo di ripetizione pari a 8.2 ms, tempo di eco pari a 3.8 ms, tempo di inversione pari a 855 ms, 240 slice sagittali e risoluzione spaziale isotropica di 1 mm. Le immagini di diffusione sono state acquisite mediante sequenza spin-echo echo-planar imaging (SE-EPI). Il protocollo prevedeva un tempo di ripetizione pari a 8400 ms, un tempo di eco pari a 85 ms, 32 direzioni di diffusione non collineari per ciascun valore di b, valori di b pari a 1000 e 2000 s/mm² e 7 immagini non pesate in diffusione, per un totale di 71 volumi. L'acquisizione comprendeva 60 slice assiali con risoluzione spaziale pari a $2.5 \times 2.5 \times 2.5$ mm³ [16]. L'acquisizione funzionale resting-state è stata eseguita mediante sequenza gradient-echo echo-planar imaging (GE-EPI). I parametri principali erano: tempo di ripetizione pari a 2400 ms, tempo di eco pari a 30 ms, 200 volumi, 40 slice assiali, risoluzione spaziale pari a $3 \times 3 \times 3$ mm³ e gap tra slice pari a 0.5 mm.

3.4 Rete Attentiva Ventrale

La Rete Attentiva Ventrale (Ventral Attention Network, VAN) è una rete funzionale coinvolta nel rilevamento di stimoli salienti, cioè stimoli che sono in grado di catturare l'attenzione anche in assenza di un orientamento volontario. Questa rete partecipa quindi al riorientamento dell'attenzione verso eventi inattesi o potenzialmente rilevanti. Nei modelli classici dell'attenzione viene distinta dalla Dorsal Attention Network, che è maggiormente coinvolta nell'orientamento volontario e top-down dell'attenzione. La VAN, invece, è più associata a meccanismi bottom-up e stimulus-driven, cioè alla capacità di interrompere un set attentivo corrente per dirigere le risorse cognitive verso uno stimolo comportamentalmente rilevante [17, 18].

Dal punto di vista anatomico-funzionale, la VAN comprende tipicamente regioni come la giunzione temporo-parietale e la corteccia frontale ventrale, con una prevalenza dell'emisfero destro [17]. Poiché queste regioni sono coinvolte in funzioni cognitive complesse, la loro dinamica può essere influenzata non solo da alterazioni corticali locali, ma anche da modificazioni delle connessioni strutturali a lungo raggio. In un modello TVB, tali connessioni sono rappresentate dalla matrice di connettività strutturale, che vincola la propagazione dell'attività tra nodi.

La scelta di focalizzarsi sulla VAN consente di ridurre i tempi computazionali e la complessità del problema rispetto a una simulazione whole-brain. L'analisi di una singola rete permette di ottimizzare i parametri biofisici in modo mirato, limitando il numero di nodi e connessioni considerate. La VAN è stata rappresentata come una sottorete composta da 20 regioni di interesse, definite sulla base della parcellazione adottata. Tali regioni sono state considerate come nodi della rete. Per ogni soggetto, le corrispondenti connessioni strutturali sono state estratte dalla matrice di connettività strutturale complessiva, ottenendo una matrice specifica per la VAN. Analogamente, i segnali BOLD associati ai nodi della rete sono stati utilizzati per ricostruire la connettività funzionale empirica e la matrice di connettività funzionale dinamica.

3.5 Pre-processing dei dati MRI

Il pre-processing dei dati di risonanza magnetica rappresenta un passaggio fondamentale del flusso di lavoro metodologico, poiché consente di ridurre il contributo di rumore, artefatti di acquisizione, movimenti del soggetto e disallineamenti spaziali tra le diverse modalità di imaging. In questo contesto, tale fase è stata particolarmente rilevante perché gli output del pre-processing costituiscono la base per la costruzione delle matrici di connettività strutturale e funzionale utilizzate successivamente nelle simulazioni implementate in The Virtual Brain.

Nel caso dei dati di diffusione, artefatti geometrici o una registrazione non accurata possono alterare la ricostruzione delle connessioni strutturali tra regioni cerebrali. Analogamente, nel caso dei dati funzionali, rumore fisiologico, movimenti del capo e drift temporali possono introdurre correlazioni spurie tra le serie temporali BOLD, influenzando la stima della connettività funzionale empirica e, di conseguenza, il confronto con la dinamica simulata.

Le procedure di pre-processing sono state applicate separatamente ai dati di diffusione e ai dati resting-state fMRI. Le immagini anatomiche T1-pesate sono state utilizzate come riferimento strutturale per le operazioni di registrazione, segmentazione e normalizzazione spaziale, in modo da garantire una corrispondenza coerente tra anatomia individuale, dati funzionali e dati di diffusione.

3.5.1 Pre-processing dei dati di diffusione

I dati di diffusione sono stati pre-processati con l'obiettivo di correggere i principali artefatti che caratterizzano le acquisizioni pesate di diffusione e di ottenere immagini adeguate alla successiva ricostruzione della connettività strutturale. La pipeline è stata basata sull'utilizzo di strumenti di MRtrix3 e della FMRIB Software Library (FSL). [19].

Come primo passaggio, è stato effettuato il denoising delle immagini di diffusione mediante il metodo MP-PCA, implementato in MRtrix3 attraverso il comando `dwidenoise`. Dopodichè si prosegue con il controllo e la correzione delle distorsioni geometriche tipiche delle sequenze echo-planar imaging, frequentemente utilizzate nelle acquisizioni di diffusione. Tali distorsioni possono derivare da disomogeneità locali del campo magnetico e risultano particolarmente evidenti in regioni cerebrali prossime a interfacce tra tessuti con diversa suscettibilità magnetica. La correzione delle distorsioni da suscettibilità è stata effettuata mediante TOPUP, che stima il campo di distorsione dovuto alle disomogeneità di suscettibilità magnetica e permette di compensare le deformazioni geometriche presenti nelle immagini.

Successivamente, è stata eseguita la correzione degli artefatti dovuti alle correnti parassite e ai movimenti del capo mediante EDDY. Le correnti parassite sono generate dalle rapide variazioni dei gradienti di diffusione e possono produrre deformazioni e disallineamenti tra i diversi volumi acquisiti. Anche piccoli movimenti del soggetto durante l'acquisizione possono compromettere la coerenza spaziale tra le immagini pesate in diffusione. La correzione combinata di questi effetti è quindi necessaria per ottenere dati più affidabili e spazialmente coerenti.

Al termine delle correzioni, l'immagine anatomica T1-pesata del singolo soggetto è stata registrata nello spazio delle immagini di diffusione. Questo è un passaggio necessario poiché la parcellazione anatomica è definita nello spa-

zio della T1, mentre la ricostruzione delle fibre e della connettività strutturale avviene nello spazio di diffusione. La trasformazione della T1, e quindi dell'atlante associato, nello spazio della diffusione permette di identificare le regioni di interesse direttamente sulle immagini utilizzate per la trattografia. In questo modo, le regioni della parcellazione possono essere utilizzate come nodi della rete per la successiva costruzione della matrice di connettività strutturale.

A partire dai dati di diffusione corretti, è stata stimata la distribuzione locale dell'orientamento delle fibre di sostanza bianca in ciascun voxel. Questa informazione è necessaria per la successiva ricostruzione trattografica, poiché consente di seguire la direzione preferenziale di diffusione dell'acqua e quindi di ricostruire i possibili percorsi delle fibre tra le regioni cerebrali considerate.

Le connessioni ottenute mediante trattografia sono state utilizzate per costruire le matrici di connettività strutturale del singolo soggetto. In particolare, sono state calcolate una matrice dei pesi, che descrive la forza della connessione tra coppie di regioni, e una matrice delle lunghezze, che rappresenta la lunghezza dei tratti ricostruiti tra le stesse regioni.

3.5.2 Pre-processing dei dati resting-state fMRI

Il pre-processing dei dati resting-state fMRI è stato eseguito con l'obiettivo di ottenere serie temporali BOLD con un miglior rapporto segnale-rumore. La pipeline è stata implementata integrando diversi software per l'elaborazione di neuroimmagini, tra cui SPM12, FSL e MRtrix3. Le principali fasi della pipeline hanno incluso: rimozione del rumore termico, slice-timing correction, realignment, coregistrazione con l'immagine anatomica, costruzione della maschera del liquido cerebrospinale, nuisance regression e filtraggio temporale.

Il primo passaggio consiste nella riduzione del rumore termico mediante Marchenko-Pastur Principal Component Analysis (MP-PCA). Questo metodo consente di separare le componenti casuali attribuibili al rumore dalle componenti informative del segnale, sfruttando le proprietà statistiche degli autovalori attesi in presenza di rumore casuale. L'immagine funzionale viene quindi ricostruita mantenendo prevalentemente le componenti associate al segnale reale. In questa pipeline, il denoising MP-PCA è stato implementato mediante il comando `dwidenoise` di MRtrix3.

Dopo la rimozione del rumore termico è stata applicata la correzione del tempo di acquisizione delle slice. Nelle acquisizioni di risonanza magnetica fun-

zionale, infatti, le diverse slice che compongono ciascun volume non vengono acquisite esattamente nello stesso istante. Questo può introdurre sfasamenti temporali tra le serie di voxel appartenenti a slice diverse. La correzione del tempo di acquisizione delle slice consente di riallineare temporalmente le slice rispetto a un tempo di riferimento comune, rendendo più coerente l'interpretazione temporale del segnale BOLD. Operativamente, il file funzionale 4D è stato suddiviso in singoli volumi 3D, corretti mediante SPM e successivamente ricombinati in un unico file 4D.

Il passaggio successivo è stato il riallineamento, finalizzato alla correzione dei movimenti del capo tra le immagini funzionali. La procedura stima trasformazioni rigide tra i diversi volumi e li riallinea a un riferimento comune, generalmente rappresentato dall'immagine funzionale media. Questo step produce sia le immagini funzionali corrette sia i parametri di movimento, che vengono successivamente utilizzati come regressori di disturbo. La correzione del movimento è particolarmente importante negli studi di connettività funzionale, poiché anche movimenti di piccola entità possono generare correlazioni artificiali tra regioni cerebrali.

Successivamente, l'immagine funzionale media è stata coregistrata all'immagine anatomica T1-pesata del soggetto. La coregistrazione consente di allineare spazialmente l'immagine funzionale all'immagine anatomica, utilizzata come immagine di riferimento, mediante una trasformazione affine, permettendo l'estrazione delle serie temporali BOLD dalle regioni di interesse definite nello spazio anatomico.

Parallelamente alla coregistrazione, è stata costruita una maschera binaria del liquido cerebrospinale per la rimozione di componenti fisiologiche non neuronali. A questo scopo, l'immagine T1-pesata è stata normalizzata allo spazio standard MNI mediante trasformazioni lineari e non lineari. L'atlante ALVIN, definito nello spazio MNI, è stato quindi trasformato nello spazio anatomico del soggetto e combinato con la segmentazione del liquido cerebrospinale ottenuta dall'immagine T1. La maschera risultante è stata sogliata con valore 0.99 ed erosa, in modo da aumentare la specificità della selezione dei voxel relativi al liquido cerebrospinale e ridurre il rischio di includere voxel appartenenti alla sostanza grigia adiacente. Infine, la maschera del liquido cerebrospinale è stata riportata nello spazio funzionale.

Una volta ottenuti i dati funzionali corretti e la maschera del liquido cerebrospinale nello spazio fMRI, è stata eseguita la nuisance regression. Questo

passaggio ha avuto lo scopo di rimuovere dal segnale BOLD componenti non direttamente attribuibili all'attività neuronale. La regressione di questi segnali consente di attenuare l'influenza di movimento, rumore fisiologico e fluttuazioni lente non informative.

Infine, al segnale corretto è stato sottoposto un filtro temporale passa-banda nell'intervallo 0.008-0.09 Hz. Questo intervallo è comunemente utilizzato negli studi resting-state per isolare le fluttuazioni lente del segnale BOLD, considerate maggiormente informative per la stima della connettività funzionale spontanea. Al termine del pre-processing, per ciascun soggetto sono state estratte le serie temporali BOLD relative alle regioni della VAN.

3.6 Connettività

Per ciascun soggetto sono state costruite matrici di connettività strutturale e funzionale statica e dinamica a partire dai dati MRI precedentemente preprocessati. Tali matrici rappresentano l'input principale per la modellizzazione cerebrali in TVB, poiché permettono di descrivere rispettivamente l'organizzazione anatomica delle connessioni tra le regioni della rete e le relazioni funzionali osservate nel segnale BOLD resting-state [10].

3.6.1 Connettività strutturale

Per ogni soggetto la matrice di connettività strutturale è stata costruita a partire dalla trattografia. Le regioni della parcellazione sono state considerate come nodi della rete, mentre le linee di flusso ricostruite tra coppie di regioni sono state utilizzate per definire gli archi. La matrice dei pesi rappresenta l'intensità della connessione strutturale tra nodi, mentre la matrice delle lunghezze contiene la lunghezza media delle linee di flusso associate a ciascuna connessione.

Prima dell'utilizzo in TVB, le matrici sono state controllate e normalizzate. La normalizzazione è necessaria perché il numero assoluto di streamlines può dipendere da fattori tecnici, come i parametri della trattografia, la dimensione delle ROI e la qualità dei dati di diffusione. L'obiettivo non è interpretare il valore assoluto delle streamlines come misura diretta del numero di fibre biologiche, ma costruire una matrice pesata che rappresenti in modo coerente la topologia e la forza relativa delle connessioni.

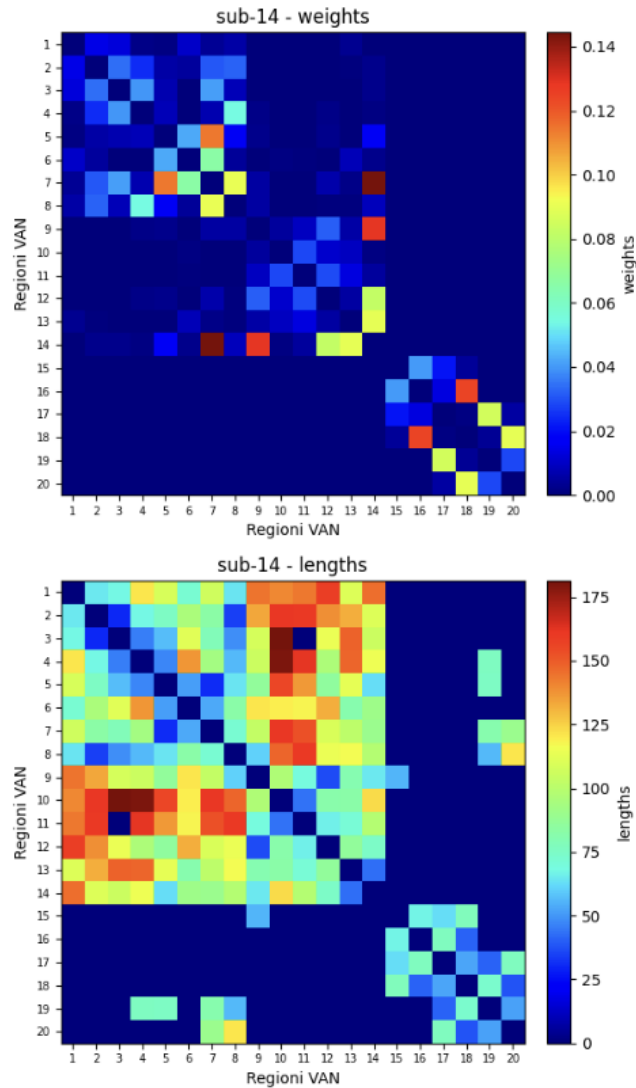


Figura 3.1: Rappresentazione della connettività strutturale del soggetto sub-14 nella Ventral Attention Network. La matrice superiore mostra i pesi delle connessioni strutturali tra le regioni della rete, mentre la matrice inferiore riporta le corrispondenti lunghezze dei tratti. Gli assi indicano le 20 regioni della VAN; ciascun elemento della matrice descrive quindi la relazione strutturale tra una coppia di regioni.

3.6.2 Connettività funzionale empirica statica

La connettività funzionale empirica (expFC) è stata stimata a partire dai segnali BOLD preprocessati dei nodi appartenenti alla VAN. Per ciascun soggetto, la serie temporale BOLD di ogni nodo della rete è stata estratta mediando

il segnale dei voxel appartenenti alla corrispondente regione di interesse definita dall'atlante. In questo modo, ciascuna regione di interesse della VAN è rappresentata da una singola serie temporale. In Figura 3.2 è riportato un esempio di segnali BOLD empirici estratti dalle diverse regioni della VAN per un singolo soggetto.

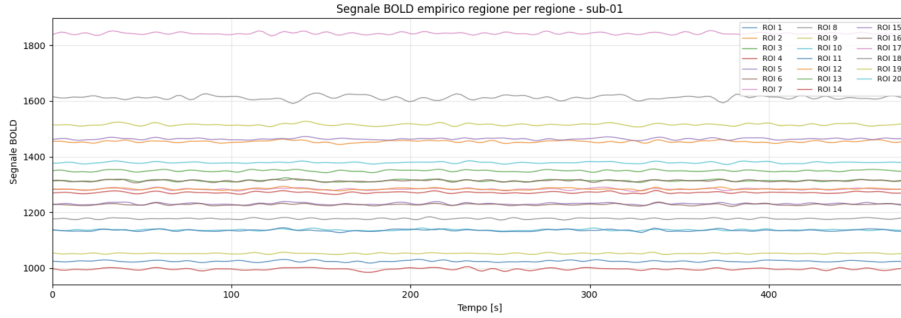


Figura 3.2: Esempio di segnali BOLD empirici estratti dalle regioni della rete attentiva ventrale per un singolo soggetto. Ogni curva rappresenta la serie temporale BOLD di una ROI della rete. L'asse delle ascisse indica il tempo in secondi, mentre l'asse delle ordinate riporta l'intensità del segnale BOLD preprocessato.

A partire da queste serie temporali, per ogni coppia di nodi è stato calcolato il coefficiente di correlazione di Pearson tra i rispettivi segnali BOLD. La matrice expFC risultante descrive quindi il grado di sincronizzazione lineare tra le fluttuazioni BOLD delle diverse regioni della rete.

Formalmente, dati due segnali BOLD $x_i(t)$ e $x_j(t)$ associati ai nodi i e j , la connettività funzionale tra i due nodi è definita come:

$$FC_{ij} = \text{corr}(x_i(t), x_j(t)) \quad (3.1)$$

La matrice expFC ottenuta tramite correlazione di Pearson è simmetrica e contiene valori compresi tra -1 e 1 . Successivamente, i valori di correlazione sono stati trasformati mediante trasformazione di Fisher-Z. Dopo la trasformazione, i valori inferiori alla soglia considerata sono stati posti a zero. Questa procedura è stata applicata per evidenziare le connessioni funzionali più rilevanti e rendere più stabile il successivo confronto tra connettività empirica e simulata. Un esempio di matrice di connettività funzionale empirica trasformata è riportato in Figura 3.3.

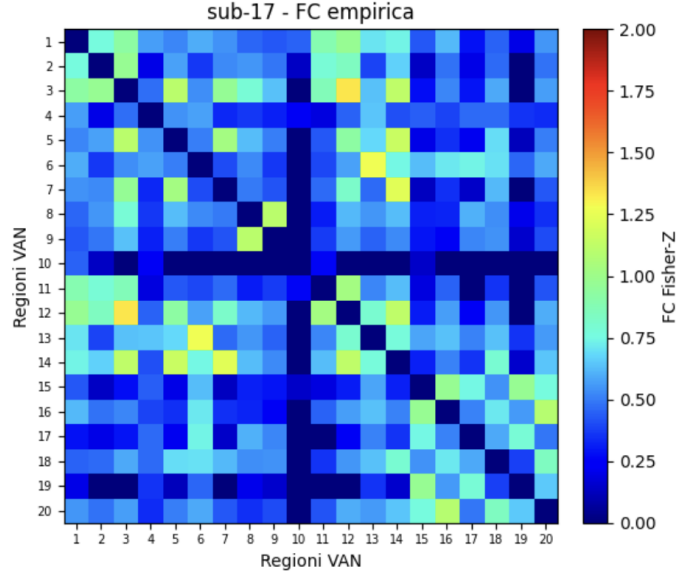


Figura 3.3: Esempio di matrice di connettività funzionale empirica trasformata relativa a un soggetto scelto randomicamente dalla coorte. La matrice è stata ottenuta calcolando il coefficiente di correlazione di Pearson tra le serie temporali BOLD delle regioni della Ventral Attention Network e applicando successivamente la trasformazione di Fisher-Z. I valori inferiori alla soglia considerata sono stati posti a zero. Le righe e le colonne rappresentano le 20 regioni della VAN, mentre la scala di colore indica l'intensità della connettività funzionale trasformata.

3.6.3 Connettività funzionale empirica dinamica

Oltre alla connettività funzionale statica, è stata considerata la connettività funzionale dinamica (Functional Connectivity Dynamics, FCD), che descrive come cambia nel tempo la connettività funzionale tra le regioni cerebrali. La FCD viene costruita calcolando la connettività funzionale su finestre temporali scorrevoli e confrontando tra loro i pattern ottenuti in finestre diverse. In questo modo, non si considera soltanto la connettività media sull'intera acquisizione, ma anche la stabilità o variabilità della configurazione funzionale nel tempo.

La FCD empirica è stata calcolata applicando finestre temporali sovrapposte della durata di 40 secondi ai segnali BOLD della rete. La finestra è stata traslata progressivamente di un tempo di ripetizione di 2.4 secondi, ottenendo una sequenza di matrici di connettività funzionale locale, ciascuna riferita a un diverso intervallo temporale dell'acquisizione. Ogni matrice di connettività

funzionale locale è stata successivamente vettorializzata e confrontata con tutte le altre mediante correlazione di Pearson. La matrice FCD risultante aveva dimensione 178×178 , dove ciascun elemento rappresenta la similarità tra due pattern di connettività funzionale calcolati in finestre temporali differenti [16].

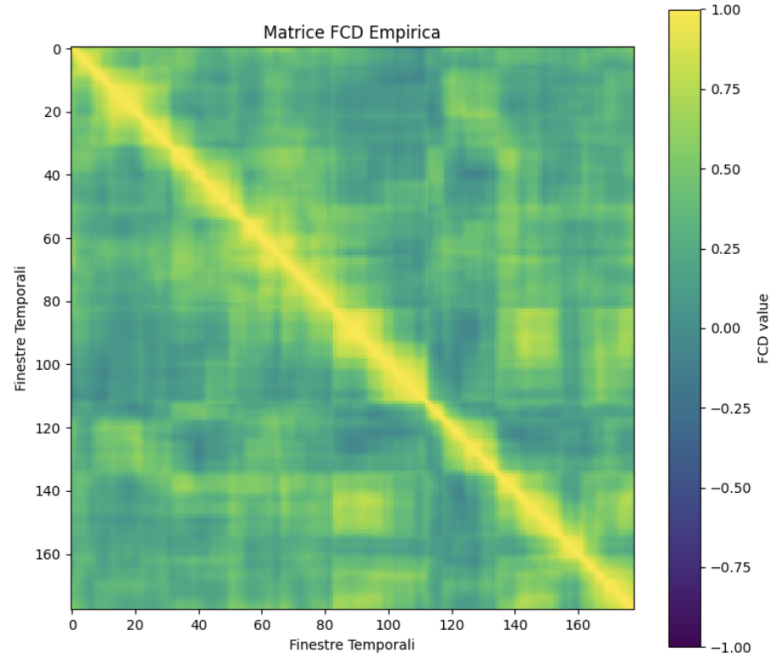


Figura 3.4: Esempio di matrice FCD empirica relativa al soggetto sub-17.

3.7 The Virtual Brain

TVB è stato utilizzato per generare segnali BOLD simulati della rete attentiva ventrale e per stimare parametri biofisici soggetto-specifici[16].

Ogni nodo della rete è stato associato a un modello dinamico locale di Wong-Wang, che descrive l'attività di popolazioni neuronali eccitatorie e inibitorie interconnesse [12, 13]. La comunicazione tra nodi è vincolata dalla matrice di connettività strutturale, mentre il segnale neurale simulato viene trasformato in segnale BOLD mediante il modello emodinamico di Balloon-Windkessel. In questo modo, il modello produce output confrontabili con i dati rs-fMRI empirici (vedi sezione 1.2).

3.7.1 Modello di Wong-Wang

Il modello utilizzato per descrivere l'attività media di popolazioni neuronali locali è il modello dinamico di Wong-Wang. In questo caso è stata utilizzata una formulazione ridotta del modello, in cui ciascun nodo della rete cerebrale è descritto da due variabili di stato che rappresentano il grado di attivazione sinaptica delle popolazioni eccitatorie e inibitorie e sono indicate con S_E e S_I . Questo modello comprende parametri fissi, mantenuti costanti durante le simulazioni, e parametri ottimizzati, che vengono modificati dall'algoritmo genetico. I parametri fissi descrivono proprietà del modello dinamico locale e sono stati impostati secondo la formulazione di Deco et al. e le indicazioni riportate nei lavori di Monteverdi et al. [13, 14, 20].

I principali parametri fissi sono riportati in Tabella 3.2.

PARAMETRI	VALORE	DESCRIZIONE
$a_e, b_e, d_e, \tau_e, W_e$	310 nC ⁻¹ , 125 Hz, 0.16 s, 100 ms, 1	Variabili di gating eccitatorie
$a_i, b_i, d_i, \tau_i, W_i$	615 nC ⁻¹ , 177 Hz, 0.087 s, 10 ms, 0.7	Variabili di gating inibitorie
γ	0.641/1000	Parametro cinetico
σ	0.01 nA	Ampiezza del rumore
I_0	0.382 nA	Input esterno effettivo complessivo

Tabella 3.2: Parametri fissi del modello di Wong-Wang. [13],[20]

I range di ricerca utilizzati dall'algoritmo sono riportati in Tabella 3.3 [20]. Tutti i parametri sono stati ottimizzati a livello soggetto-specifico.

PARAMETRO	LIMITE INFERIORE	LIMITE SUPERIORE	DESCRIZIONE
G	0.1	4.0	Fattore di scala dell'accoppiamento globale
J_{NMDA}	0.115	0.19	Accoppiamento sinaptico eccitatorio
J_i	0.5	3.5	Accoppiamento sinaptico inibitorio
w_+	1.1	1.85	Ricorrenza eccitatoria locale

Tabella 3.3: Parametri ottimizzati tramite NSGA-II e relativi intervalli di ricerca.[20]

I parametri di acquisizione con cui le simulazioni sono state impostate sono il passo di integrazione temporale $dt = 0.1$ ms, la velocità di conduzione pari a 12.5 m/s, ovvero la velocità di propagazione del segnale neurale tra regioni cerebrali, la durata simulata pari a 480000 ms e il tempo di ripetizione pari a 2400 ms. I primi quattro volumi della simulazione sono stati esclusi dall'analisi per ridurre l'influenza del transitorio iniziale. Le simulazioni hanno prodotto segnali BOLD simulati, dai quali sono state ricostruite simFC e simFCD con procedure analoghe a quelle applicate ai dati empirici.

3.7.2 Simulazione dell'attività della Rete Attentiva Ventrle

Per ciascun soggetto e per ogni combinazione parametrica valutata, TVB ha prodotto un segnale BOLD simulato per le regioni appartenenti alla VAN. Per ridurre l'influenza del transitorio iniziale i primi quattro volumi della simulazione sono stati esclusi dall'analisi. A partire dai segnali BOLD simulati è stata calcolata la matrice di connettività funzionale simulata mediante correlazione di Pearson tra le serie temporali delle diverse regioni della rete. In modo analogo a quanto effettuato per i dati empirici, la matrice è stata successivamente trasformata mediante Fisher-Z e sogliata (vedi sezione 3.6.2).

Nel confronto con le simulazioni, la parte triangolare superiore della matrice, esclusa la diagonale, è stata vettorializzata e utilizzata per calcolare il coefficiente di correlazione di Pearson tra connettività empirica e simulata. Questo approccio consente di evitare la duplicazione delle informazioni

contenute nella matrice simmetrica e permette di confrontare direttamente le connessioni corrispondenti tra le due matrici.

Dagli stessi segnali simulati è stata inoltre ricostruita la matrice FCD simulata. Anche in questo caso è stata applicata la stessa procedura descritta per i dati empirici: suddivisione del segnale BOLD in finestre temporali scorrevoli, calcolo della connettività funzionale locale per ciascuna finestra e successivo confronto tra i pattern funzionali ottenuti nei diversi istanti temporali.

In questo modo, per ogni simulazione sono stati ottenuti output direttamente confrontabili con le corrispondenti grandezze empiriche.

3.8 Ottimizzazione multi-parametrica

L'obiettivo dell'ottimizzazione è individuare per ciascun soggetto la combinazione di parametri del modello in grado di minimizzare le differenze tra la dinamica simulata e dati empirici. Come metodo standard è stata considerata una procedura basata su grid search, in cui i parametri vengono analizzati a coppie. A partire da questo approccio, sono stati proposti due nuovi metodi, diversi tra loro per gli obiettivi considerati, ma entrambi basati sull'algoritmo genetico multi-obiettivo NSGA-II. Il primo metodo valuta la qualità della simulazione attraverso il confronto tra connettività funzionale statica e dinamica simulate ed empiriche, utilizzando come obiettivi $1 - PCC$ e KS , proprio come il metodo standard. Il secondo metodo si basa su obiettivi feature-based, definiti a partire da proprietà sintetiche estratte dal segnale BOLD simulato e dalla rete funzionale.

3.8.1 Grid-search

Come approccio di riferimento è stata considerata una procedura di ottimizzazione mediante *grid-search*. In questo caso, i parametri biofisici non sono stati esplorati simultaneamente nello spazio a quattro dimensioni, ma ottimizzati a coppie, mantenendo fissati gli altri parametri [20]. Nello specifico vengono analizzati per coppia G con J_i e J_N con w_p . Per ciascuna coppia viene definita una griglia di valori possibili e, per ogni combinazione, viene eseguita una simulazione del modello.

La funzione di costo da minimizzare è definita come la somma di $1 - PCC$ e della distanza KS . Il primo termine descrive la dissimilarità tra la matrice

di connettività funzionale statica simulata e quella empirica, mentre il secondo misura la differenza tra le distribuzioni di FCD simulata ed empirica:

$$Cost = (1 - PCC(simFC, expFC)) + KS(simFCD, expFCD). \quad (3.2)$$

3.8.2 NSGA-II e definizione degli obiettivi

L'ottimizzazione dei parametri TVB è stata implementata in Python utilizzando il pacchetto `pymoo`, che fornisce strumenti per la definizione di problemi di ottimizzazione e l'applicazione di algoritmi evolutivi multi-obiettivo [21].

Per ciascun soggetto, ogni individuo della popolazione veniva rappresentato con una diversa combinazione dei quattro parametri e a ogni iterazione, questa combinazione proposta veniva utilizzata per eseguire una simulazione TVB. Successivamente, il segnale BOLD simulato veniva confrontato con il dato empirico mediante le funzioni obiettivo.

Le soluzioni vengono ordinate in fronti di Pareto mediante la procedura di non-dominated sorting, e allo stesso tempo viene mantenuta la diversità delle soluzioni all'interno della popolazione attraverso la crowding distance.

L'ottimizzazione è stata eseguita con popolazione pari a 30 individui e 40 generazioni. Per ogni soggetto e per ciascuna strategia di ottimizzazione, l'algoritmo ha prodotto un insieme di soluzioni non dominate, corrispondente a un fronte di Pareto soggetto-specifico. Ogni punto del fronte rappresenta una diversa combinazione dei parametri e identifica un possibile compromesso tra gli obiettivi considerati. Al termine dell'ottimizzazione è stato selezionato per ciascun soggetto un singolo punto, considerato rappresentativo del compromesso tra gli obiettivi e utilizzato nelle analisi successive.

Ottimizzazione basata su PCC e KS

In questo metodo, la bontà della simulazione è stata valutata confrontando la matrice di connettività funzionale simulata ($simFC$) con la matrice di connettività funzionale empirica ($expFC$) e la matrice di connettività funzionale dinamica simulata ($simFCD$) con quella empirica ($expFCD$). Sono state definite due funzioni obiettivo: la minimizzazione di $1 - PCC$ tra $simFC$ ed $expFC$ e la minimizzazione della distanza di Kolmogorov-Smirnov (KS) tra $simFCD$ ed $expFCD$.

$$f_1 = 1 - PCC(simFC, expFC). \quad (3.3)$$

In questo modo, valori più bassi di f_1 indicano una maggiore somiglianza tra la connettività funzionale simulata e quella empirica. Il secondo obiettivo è definito come:

$$f_2 = KS(simFCD, expFCD). \quad (3.4)$$

La distanza KS misura la differenza tra le distribuzioni dei valori della FCD simulata ed empirica. Anche in questo caso, valori più bassi indicano una maggiore somiglianza tra la dinamica funzionale simulata e quella osservata. L'ottimizzazione cerca quindi soluzioni che minimizzino simultaneamente $1 - PCC$ e KS .

Ottimizzazione *feature-based*

Prima di definire la strategia di ottimizzazione è stata effettuata una fase preliminare di selezione delle feature. L'obiettivo era individuare misure che fossero maggiormente informative per descrivere le differenze tra i gruppi della coorte. La selezione è stata effettuata mediante il software Orange, utilizzando un criterio di ranking basato sull'*Information Gain*. Questa misura consente di valutare quanto una determinata feature contribuisca a ridurre l'incertezza nella distinzione tra i gruppi considerati. Di conseguenza, valori più elevati di *Information Gain* indicano feature potenzialmente più rilevanti e discriminative. Sono state considerate misure derivate dal segnale BOLD, dalla connettività funzionale statica e dalla connettività funzionale dinamica. Tra queste, sono state selezionate per l'ottimizzazione le feature con maggiore capacità discriminativa, e sono il centroide spettrale del segnale BOLD, il valore medio della FCD, i nodi core e la densità della connettività funzionale.

Il centroide spettrale è stato utilizzato per sintetizzare il contenuto in frequenza del segnale BOLD. Questa misura può essere interpretata come una frequenza media pesata sullo spettro di potenza: valori più alti indicano che una quota maggiore della potenza del segnale è distribuita verso frequenze relativamente elevate, mentre valori più bassi sono associati a dinamiche più lente. Questa feature consente di valutare se la simulazione è in grado di riprodurre non solo l'ampiezza del segnale BOLD, ma anche la sua organizzazione nel dominio della frequenza.

Il valore medio della FCD (indicato come FCDmean) descrive il grado medio di somiglianza tra configurazioni di connettività funzionale ottenute in finestre temporali differenti. Valori elevati di questa feature indicano che i pattern di connettività tendono a rimanere simili nel tempo, suggerendo una dinamica funzionale più stabile, mentre valori più bassi riflettono una maggiore variabilità temporale, cioè una più marcata riorganizzazione delle configurazioni funzionali durante la simulazione.

Infine, sono state considerate feature derivate dalla teoria dei grafi, in particolare i nodi core e la densità. I nodi core identificano le regioni che presentano un ruolo maggiormente centrale nella rete, in quanto caratterizzate da un più alto livello di connessione rispetto agli altri nodi. Questa misura permette di descrivere la presenza di un nucleo funzionale più integrato. La densità della rete, invece, quantifica la proporzione di connessioni effettivamente presenti, dopo l'applicazione di una soglia, rispetto al numero massimo di connessioni possibili. Una densità più elevata indica una rete più compatta e interconnessa, mentre una densità più bassa suggerisce una configurazione più sparsa.

Queste feature sono state utilizzate per costruire le funzioni obiettivo dell'ottimizzazione *feature-based*, confrontando i valori ottenuti dalle simulazioni con quelli calcolati sui dati empirici. Sono state descritte tre funzioni obiettivo:

$$f_1 = RMSE(centroide_{sim}, centroide_{emp}), \quad (3.5)$$

$$f_2 = RMSE(FCDmean_{sim}, FCDmean_{emp}), \quad (3.6)$$

$$f_3 = \sqrt{(nodi\ core_{sim} - nodi\ core_{emp})^2 + (densita_{sim} - densita_{emp})^2}. \quad (3.7)$$

Il terzo obiettivo integra due proprietà della rete funzionale. Questa scelta consente di includere informazioni sull'organizzazione topologica della rete senza aumentare il numero complessivo di funzioni obiettivo. In questo modo, la strategia feature-based rimane formulata come un problema a tre obiettivi, evitando di rendere eccessivamente complesso il fronte di Pareto. Rispetto alla strategia basata su PCC e KS, che valuta la somiglianza globale tra matrici simulate ed empiriche, l'approccio feature-based è quindi maggiormente orientato alla riproduzione di proprietà specifiche della dinamica BOLD e

dell'organizzazione funzionale della rete.

3.8.3 Selezione del punto ottimo

Criterio adottato

Per entrambe le strategie di ottimizzazione è stato scelto il criterio della distanza. Tale criterio è stato applicato sia al fronte bidimensionale ottenuto con il metodo basato su PCC e KS, sia al fronte tridimensionale ottenuto con il metodo feature-based. La scelta è stata motivata dalla necessità di mantenere una procedura di selezione uniforme tra i due approcci, così da rendere più diretto e coerente il confronto dei parametri ottimizzati.

Prima del calcolo della distanza, i valori degli obiettivi sono stati normalizzati mediante normalizzazione min-max, in modo da rendere confrontabili metriche espresse su scale differenti. Successivamente, per ciascuna soluzione del fronte è stata calcolata la distanza euclidea rispetto al punto ideale.

Valutazione a posteriori delle strategie di ottimizzazione

Al termine dell'ottimizzazione, per ogni soggetto e per ciascuna strategia è stata selezionata una singola soluzione rappresentativa dal rispettivo fronte di Pareto. A partire da questa soluzione finale, è stata poi effettuata una valutazione a posteriori, con l'obiettivo di confrontare i due metodi anche rispetto alle metriche non utilizzate direttamente nella fase di ottimizzazione.

In particolare, nel caso delle soluzioni ottenute tramite ottimizzazione feature-based, sono stati calcolati successivamente anche i valori di $1 - \text{PCC}$ e KS . In questo modo è stato possibile valutare quanto le soluzioni individuate sulla base delle feature sintetiche fossero comunque in grado di riprodurre la connettività funzionale statica e dinamica. Allo stesso modo, per le soluzioni ottenute con il metodo basato sull'ottimizzazione di PCC e KS, sono stati calcolati a posteriori gli obiettivi feature-based, così da verificare il comportamento di tali soluzioni rispetto alle proprietà sintetiche del segnale BOLD e della rete funzionale.

Questa procedura ha permesso di confrontare le due strategie su un insieme comune di metriche. Poiché i due metodi sono stati ottimizzati rispetto a obiettivi differenti, il confronto dei soli valori usati durante l'ottimizzazione non sarebbe stato sufficiente per descrivere in modo completo le loro presta-

zioni. La valutazione a posteriori consente quindi di verificare se una soluzione selezionata da un metodo mantenga buoni risultati anche rispetto ai criteri considerati dall'altro approccio.

Capitolo 4

Risultati

In questo capitolo vengono riportati i risultati ottenuti dall'ottimizzazione multi-obiettivo del modello TVB. Per ogni soggetto è stato ottenuto un fronte di Pareto, formato da soluzioni non dominate corrispondenti a diverse combinazioni dei parametri del modello. A partire da ciascun fronte è stata poi selezionata una singola soluzione rappresentativa, indicata come punto ottimo, così da poter confrontare i risultati non solo tra soggetti, ma anche tra gruppi clinici e tra strategie di ottimizzazione.

Inizialmente viene analizzato il confronto tra l'approccio basato su grid search e l'ottimizzazione mediante NSGA-II basata sul calcolo di PCC e KS. Questo confronto permette di valutare se l'utilizzo dell'algoritmo genetico consenta di ottenere soluzioni più informative rispetto al metodo tradizionale. Si procede con la valutazione tra ottimizzazione basata sul calcolo di PCC e KS e ottimizzazione feature-based, dove le feature estratte sono state scelte tramite un criterio di selezione basato sull'Information Gain. Dopodichè si analizza la distribuzione dei parametri biofisici ottimizzati nei tre gruppi clinici, i punti ottimi ottenuti con il criterio di selezione e i fronti di Pareto da cui tali soluzioni derivano, con l'obiettivo di evidenziare il compromesso tra gli obiettivi dell'ottimizzazione.

4.1 Confronto tra metodi di ottimizzazione

4.1.1 Confronto tra grid search e NSGA-II basato su PCC e KS

Come primo confronto tra metodi di ottimizzazione si è deciso di andare a valutare le prestazioni della grid search e dell'algoritmo NSGA-II basato su PCC e KS. Gli obiettivi confrontati sono $1 - PCC$, relativo al confronto tra la matrice di connettività funzionale statica empirica e simulata, e KS , relativo al confronto tra le matrici FCD empirica e simulata. I risultati relativi a questo confronto sono riportati in Figura 4.1 , Figura 4.2 e Figura 4.3 .

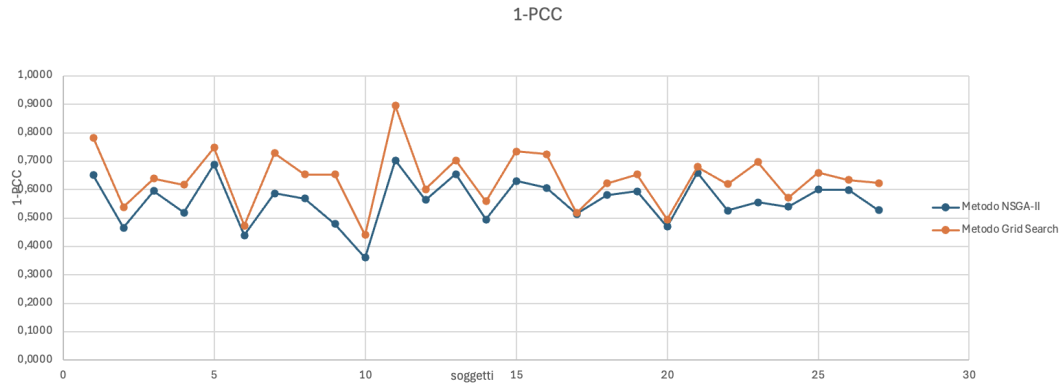


Figura 4.1: Confronto sui valori di 1-PCC tra metodo di ottimizzazione utilizzando l'algoritmo genetico NSGA-II basato su PCC e KS e metodo di ottimizzazione utilizzando la grid search. La linea blu indica il metodo NSGA-II basato su PCC e KS, mentre la linea arancione indica la grid search. Valori inferiori di $1 - PCC$ corrispondono a una maggiore somiglianza tra la connettività funzionale statica empirica e simulata.

Dal confronto soggetto per soggetto emerge che NSGA-II presenta valori di $1 - PCC$ inferiori rispetto alla grid search in tutti i soggetti analizzati.

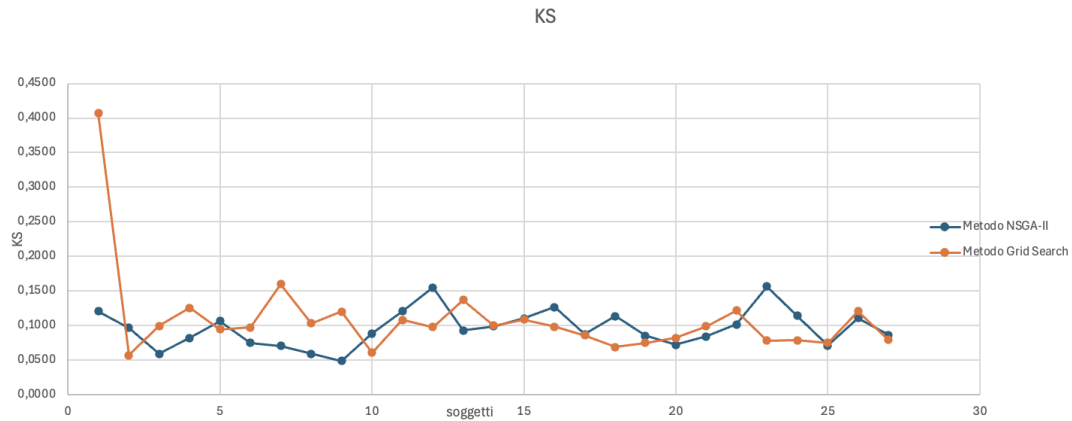


Figura 4.2: Confronto sui valori di KS tra metodo di ottimizzazione utilizzando l'algoritmo genetico NSGA-II e metodo di ottimizzazione utilizzando la grid search. La linea blu indica il metodo NSGA-II basato su PCC e KS , mentre la linea arancione indica la grid search. Valori inferiori di KS corrispondono a una maggiore somiglianza tra le distribuzioni di FCD empirica e simulata.

Per quanto riguarda KS , il confronto soggetto per soggetto mostra un andamento più variabile rispetto a quanto osservato per $1 - PCC$. In 14 soggetti su un totale di 27 il metodo NSGA-II presenta valori inferiori rispetto alla grid search, mentre nei restanti 13 soggetti la grid search mostra valori più bassi. È osservabile un valore molto più elevato degli altri per la grid search al primo soggetto, che contribuisce all'aumento della variabilità di questo obiettivo.

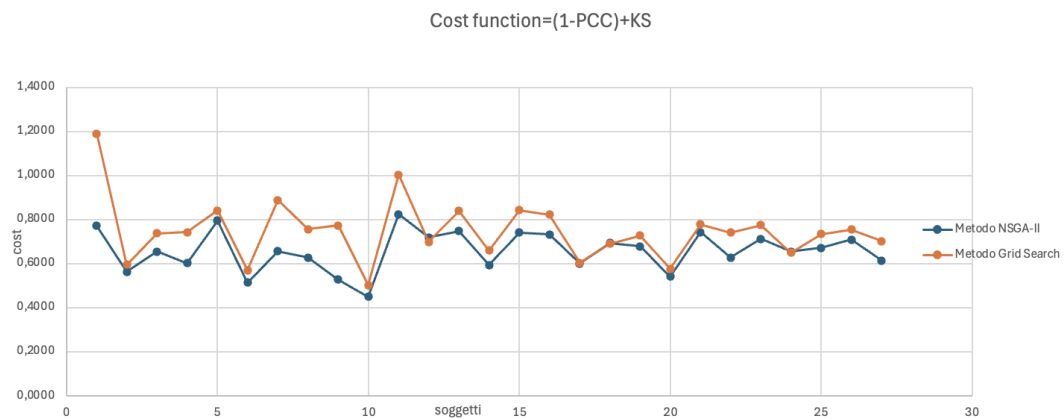


Figura 4.3: Confronto tra i metodi di ottimizzazione sui valori della somma tra i due obiettivi da minimizzare. La linea blu indica il metodo NSGA-II basato su PCC e KS , mentre la linea arancione indica la grid search.

Si è analizzata anche la funzione di costo complessiva, calcolata come somma degli obiettivi $(1 - PCC) + KS$. Questa funzione mostra valori inferiori

per NSGA-II rispetto alla grid search ed è osservabile in 24 su 27 soggetti analizzati.

Per valutare la presenza di differenze significative tra i metodi grid search e NSGA-II, è stato applicato il test dei ranghi con segno di Wilcoxon per dati appaiati. Il test ha evidenziato una differenza statisticamente significativa per l'obiettivo 1 – PCC e per la funzione di costo complessiva, con valori inferiori nel caso di NSGA-II. Questo indica che NSGA-II tende a migliorare soprattutto la somiglianza tra connettività funzionale empirica e simulata. Per l'obiettivo KS , invece, non è stata osservata una differenza statisticamente significativa, suggerendo che il miglioramento relativo alla FCD è meno marcato.

4.1.2 Confronto tra metodo basato su PCC e KS e metodo feature-based

Un altro confronto che è stato fatto è quello di andare ad analizzare gli obiettivi tra il metodo di ottimizzazione basato su matrici di connettività funzionale statica e dinamica, ovvero il metodo PCC e KS, e il metodo basato sul calcolo di feature rappresentative. Per confrontare i due metodi è stata utilizzata la stessa logica di analisi. Per ciascun metodo è stato prima selezionato il punto ottimo sul fronte di Pareto. Successivamente, per lo stesso punto, sono stati calcolati anche gli obiettivi dell'altro metodo. In questo modo, ogni soluzione ottima è stata descritta da cinque misure complessive: 1-PCC, KS, RMSE del centroide spettrale, RMSE del valore medio di FCD e RMSE complessivo di nodi core e densità. I risultati relativi a questo confronto sono riportati in Figura 4.4 , Figura 4.5, Figura 4.6, Figura 4.7 e Figura 4.8.

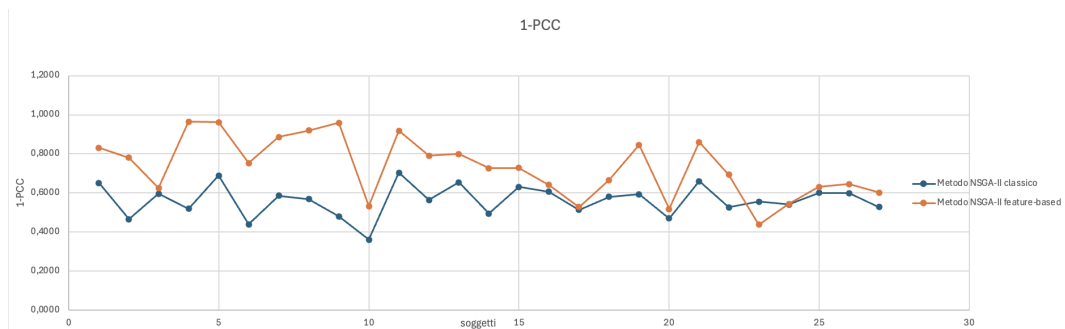


Figura 4.4: Confronto sui valori di 1-PCC tra metodo di ottimizzazione utilizzando l'algoritmo genetico NSGA-II basato su PCC e KS e metodo feature-based.

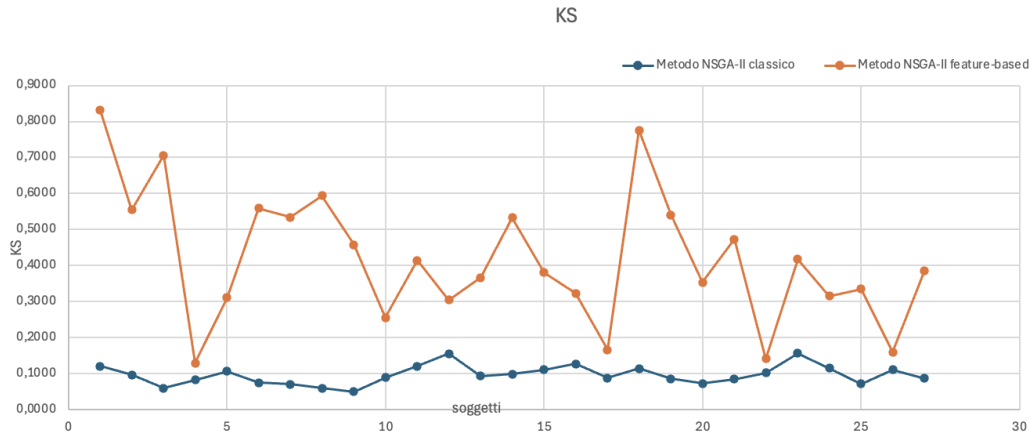


Figura 4.5: Confronto sulla distanza KS tra metodo di ottimizzazione utilizzando l'algoritmo genetico NSGA-II basato su PCC e KS e metodo feature-based.

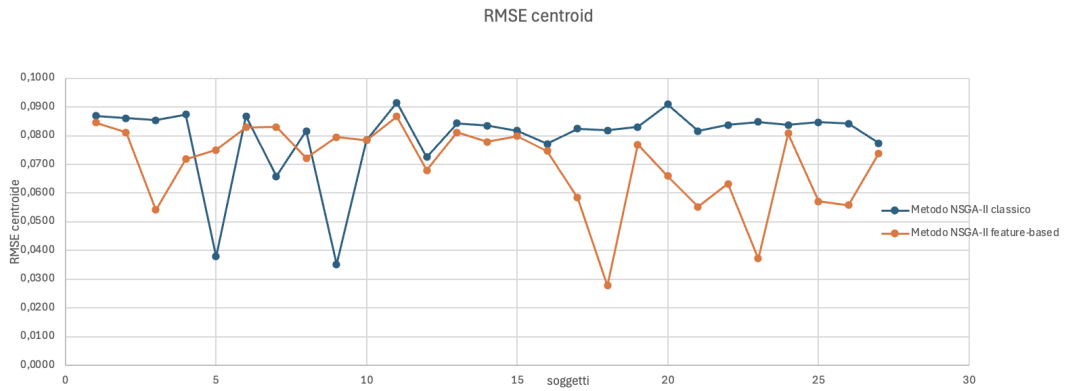


Figura 4.6: RMSE del centroide spettrale del metodo di ottimizzazione utilizzando l'algoritmo genetico NSGA-II basato su PCC e KS e del metodo feature-based.

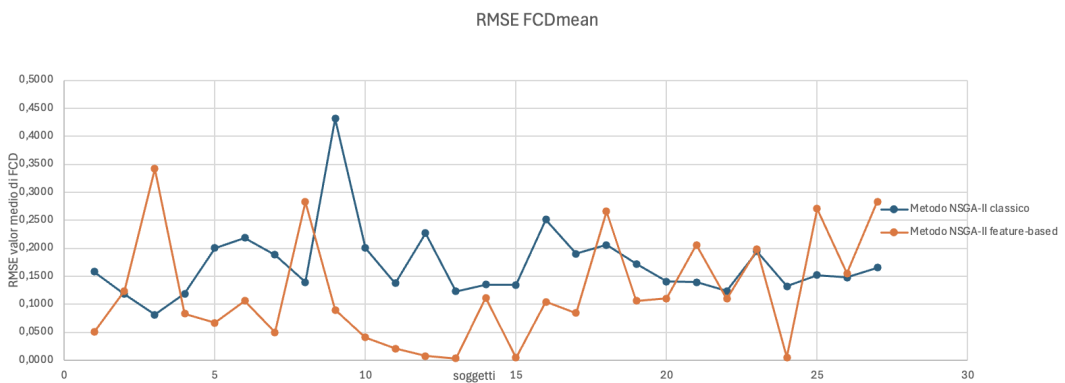


Figura 4.7: RMSE del valor medio di FCD del metodo di ottimizzazione utilizzando l'algoritmo genetico NSGA-II basato su PCC e KS e del metodo feature-based.

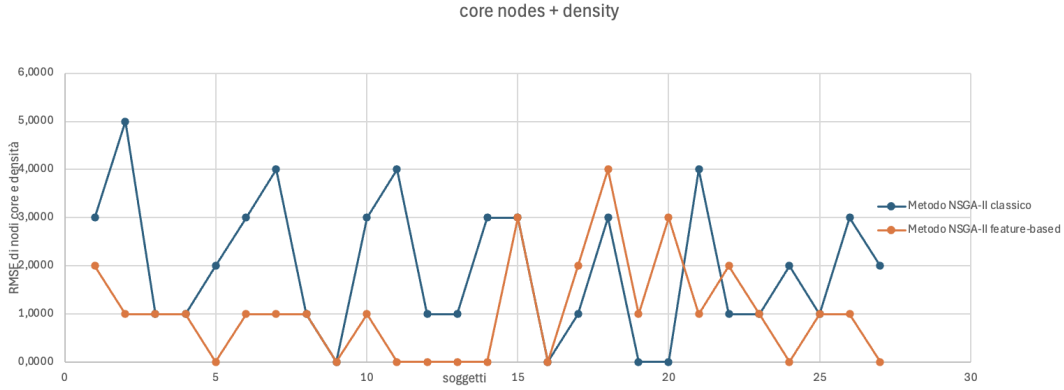


Figura 4.8: Confronto della metrica costituita da nodi core e densità tra metodo di ottimizzazione utilizzando l'algoritmo genetico NSGA-II basato su PCC e KS e del metodo feature-based.

Dai grafici emerge che ciascun metodo tende a ottenere risultati migliori soprattutto sugli obiettivi che minimizza direttamente. Per esempio, il metodo basato su PCC e KS mostra valori più bassi di 1-PCC per tutti i soggetti tranne uno, e valori più bassi di KS in tutti i soggetti. Al contrario, il metodo feature-based risulta più vantaggioso per gli obiettivi legati alle feature, poiché queste misure entrano direttamente nella funzione obiettivo. Questi risultati indicano che i due metodi tendono a privilegiare gli obiettivi inclusi direttamente nella rispettiva procedura di ottimizzazione. Questo comportamento è stato confermato anche dai test statistici effettuati mediante il test dei ranghi con segno di Wilcoxon.

4.2 Distribuzione dei parametri ottimizzati

Dopo aver confrontato i vari metodi di ottimizzazione è stata analizzata la distribuzione dei parametri biofisici all'interno delle soluzioni che costituiscono il fronte di Pareto soggetto-specifico. Questa analisi permette di osservare come i diversi parametri del modello si distribuiscano e di valutare, in modo preliminare, eventuali differenze nella variabilità dei parametri tra soggetti, gruppi clinici e strategie di ottimizzazione.

Sono stati considerati i quattro parametri ottimizzati del modello di Wong-Wang, ovvero l'accoppiamento globale G , il parametro di inibizione locale J_i , la forza della trasmissione eccitatoria mediata da recettori NMDA J_N e il peso della ricorrenza eccitatoria locale w_p . I valori assunti da questi parametri nelle soluzioni appartenenti al fronte di Pareto sono stati rappresentati mediante

boxplot, separando i soggetti nei tre gruppi clinici. La distribuzione dei valori ottenuti con il metodo basato su PCC e KS è mostrata in Figura 4.9, 4.10, 4.11 e 4.12, mentre per la distribuzione dei valori ottenuti con il metodo feature-based è mostrata in Figura 4.13, 4.14, 4.15 e 4.16.

4.2.1 Metodo basato su PCC e KS

Parametro G

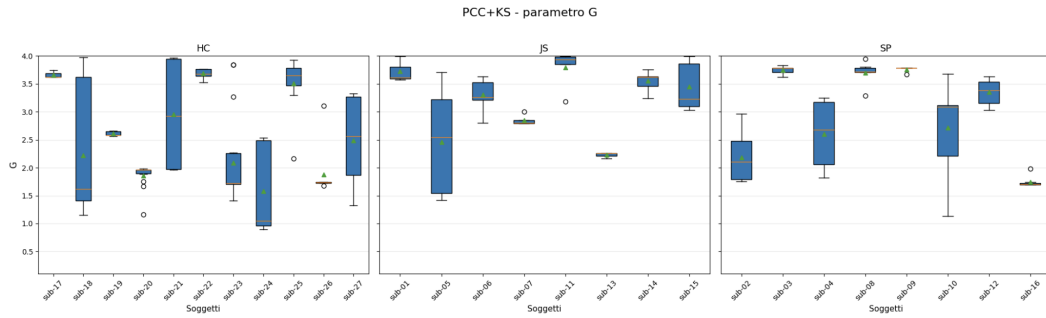


Figura 4.9: Boxplot del parametro G per tutti i soggetti suddivisi nei 3 gruppi clinici. La figura mostra la distribuzione dei valori di G permettendo di osservare la variabilità intra-soggetto e le differenze tra gruppi. La linea arancio interna indica la mediana, il triangolo verde indica la media, mentre i limiti inferiore e superiore del box rappresentano rispettivamente il primo e il terzo quartile. I baffi mostrano l'estensione dei valori osservati, mentre eventuali punti esterni indicano la presenza di outlier.

Parametro J_i

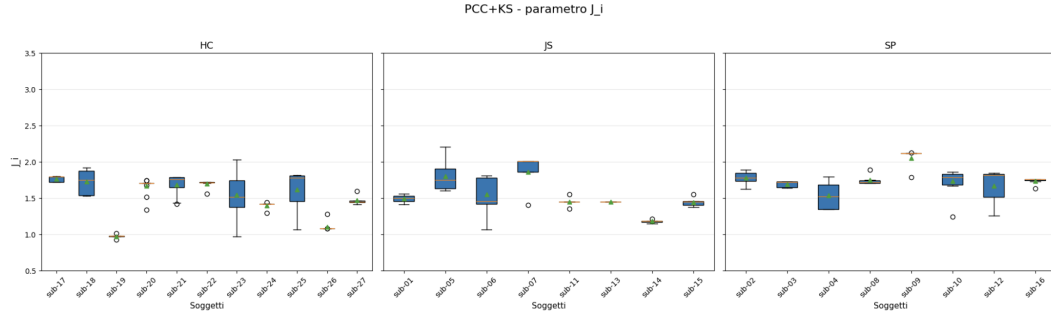


Figura 4.10: Boxplot del parametro J_i per tutti i soggetti suddivisi nei 3 gruppi clinici. La figura mostra la distribuzione dei valori di J_i permettendo di osservare la variabilità intra-soggetto e le differenze tra gruppi. La linea arancio interna indica la mediana, il triangolo verde indica la media, mentre i limiti inferiore e superiore del box rappresentano rispettivamente il primo e il terzo quartile. I baffi mostrano l'estensione dei valori osservati, mentre eventuali punti esterni indicano la presenza di outlier.

Parametro J_N

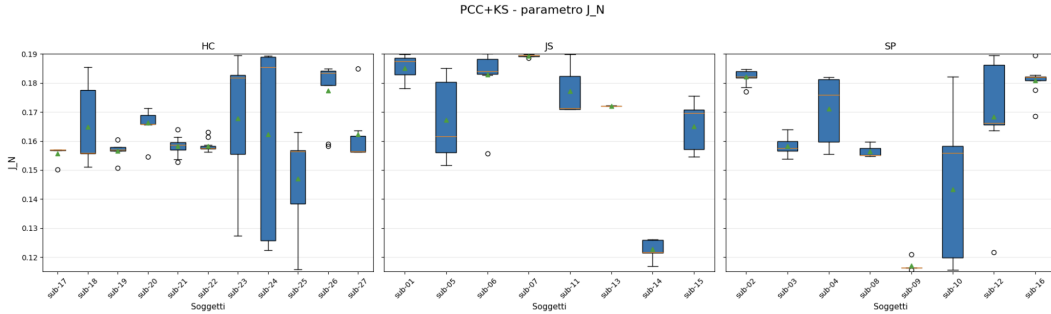


Figura 4.11: Boxplot del parametro J_N per tutti i soggetti suddivisi nei 3 gruppi clinici. La figura mostra la distribuzione dei valori di J_N permettendo di osservare la variabilità intra-soggetto e le differenze tra gruppi. La linea arancio interna indica la mediana, il triangolo verde indica la media, mentre i limiti inferiore e superiore del box rappresentano rispettivamente il primo e il terzo quartile. I baffi mostrano l'estensione dei valori osservati, mentre eventuali punti esterni indicano la presenza di outlier.

Parametro w_p

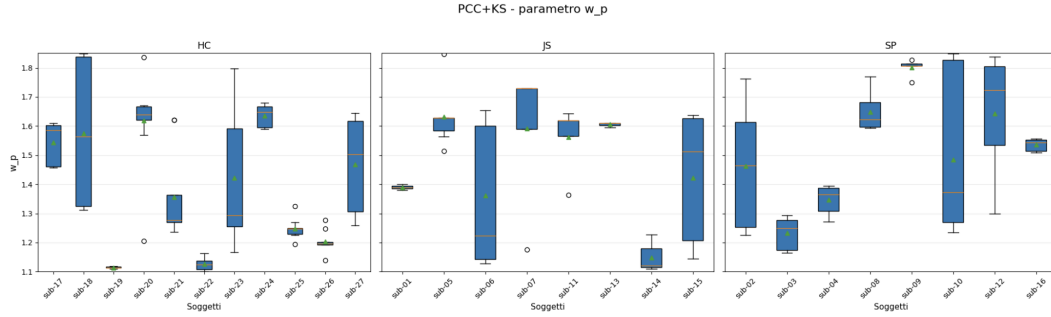


Figura 4.12: Boxplot del parametro w_p per tutti i soggetti suddivisi nei 3 gruppi clinici. La figura mostra la distribuzione dei valori di w_p permettendo di osservare la variabilità intra-soggetto e le differenze tra gruppi. La linea arancio interna indica la mediana, il triangolo verde indica la media, mentre i limiti inferiore e superiore del box rappresentano rispettivamente il primo e il terzo quartile. I baffi mostrano l'estensione dei valori osservati, mentre eventuali punti esterni indicano la presenza di outlier.

4.2.2 Metodo feature-based

Parametro G

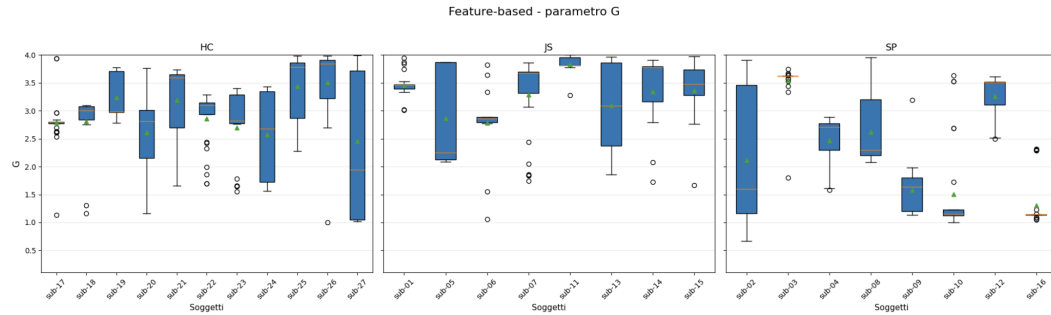


Figura 4.13: Boxplot del parametro G per tutti i soggetti suddivisi nei 3 gruppi clinici. La figura mostra la distribuzione dei valori di G permettendo di osservare la variabilità intra-soggetto e le differenze tra gruppi. La linea arancio interna indica la mediana, il triangolo verde indica la media, mentre i limiti inferiore e superiore del box rappresentano rispettivamente il primo e il terzo quartile. I baffi mostrano l'estensione dei valori osservati, mentre eventuali punti esterni indicano la presenza di outlier.

Parametro J_i

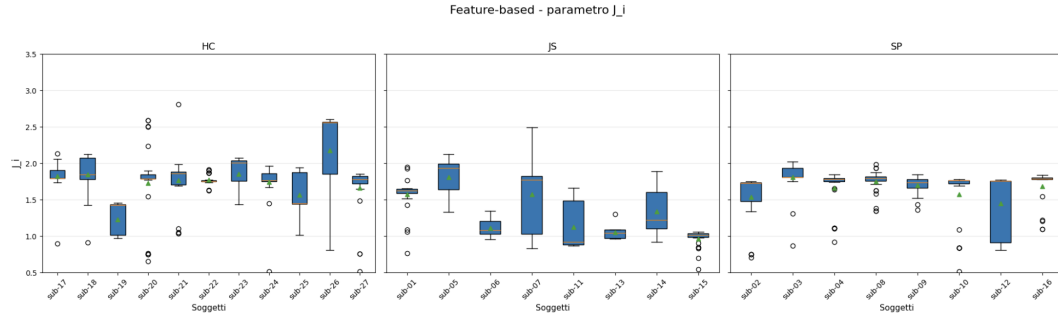


Figura 4.14: Boxplot del parametro J_i per tutti i soggetti suddivisi nei 3 gruppi clinici. La figura mostra la distribuzione dei valori di J_i permettendo di osservare la variabilità intra-soggetto e le differenze tra gruppi. La linea arancio interna indica la mediana, il triangolo verde indica la media, mentre i limiti inferiore e superiore del box rappresentano rispettivamente il primo e il terzo quartile. I baffi mostrano l'estensione dei valori osservati, mentre eventuali punti esterni indicano la presenza di outlier.

Parametro J_N

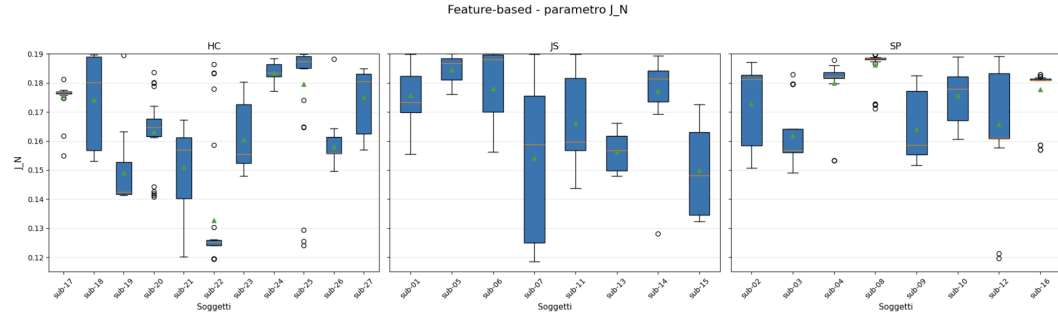


Figura 4.15: Boxplot del parametro J_N per tutti i soggetti suddivisi nei 3 gruppi clinici. La figura mostra la distribuzione dei valori di J_N permettendo di osservare la variabilità intra-soggetto e le differenze tra gruppi. La linea arancio interna indica la mediana, il triangolo verde indica la media, mentre i limiti inferiore e superiore del box rappresentano rispettivamente il primo e il terzo quartile. I baffi mostrano l'estensione dei valori osservati, mentre eventuali punti esterni indicano la presenza di outlier.

Parametro w_p

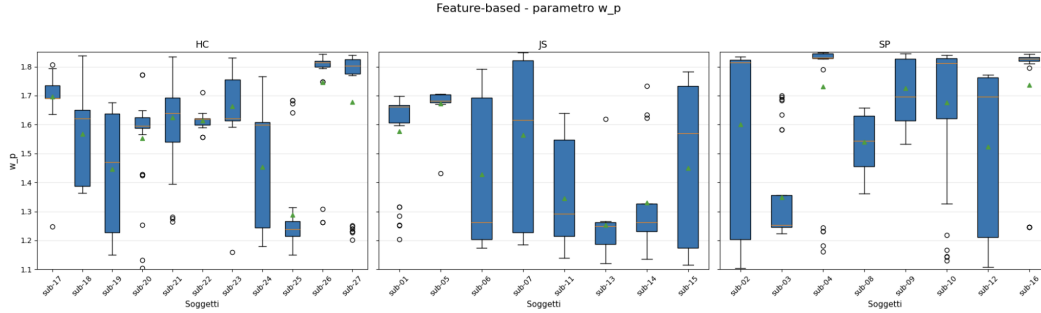


Figura 4.16: Boxplot del parametro w_p per tutti i soggetti suddivisi nei 3 gruppi clinici. La figura mostra la distribuzione dei valori di w_p permettendo di osservare la variabilità intra-soggetto e le differenze tra gruppi. La linea arancio interna indica la mediana, il triangolo verde indica la media, mentre i limiti inferiore e superiore del box rappresentano rispettivamente il primo e il terzo quartile. I baffi mostrano l'estensione dei valori osservati, mentre eventuali punti esterni indicano la presenza di outlier.

4.2.3 Fronte di Pareto

In parallelo alla visualizzazione della distribuzione dei parametri, è utile fornire una rappresentazione visiva del processo di ottimizzazione. Per mostrare in modo qualitativo il comportamento delle due strategie di ottimizzazione, sono riportati di seguito alcuni esempi di fronti di Pareto relativi sia al metodo basato su PCC e KS sia al metodo feature-based, riportati in Figura 4.17 e Figura 4.19. Nel caso del metodo basato su PCC e KS il fronte di Pareto è bi-dimensionale, perché l'ottimizzazione considera due funzioni obiettivo, mentre nel metodo feature-based il fronte di Pareto è tridimensionale, perché l'ottimizzazione viene guidata da tre obiettivi. L'obiettivo è fornire un esempio visivo di come le soluzioni non dominate si distribuiscono nello spazio degli obiettivi e di come viene selezionato il punto ottimo. La selezione del punto ottimo all'interno del fronte di Pareto è mostrata in Figura 4.18 e in Figura 4.20.

Ottimizzazione basata su PCC e KS

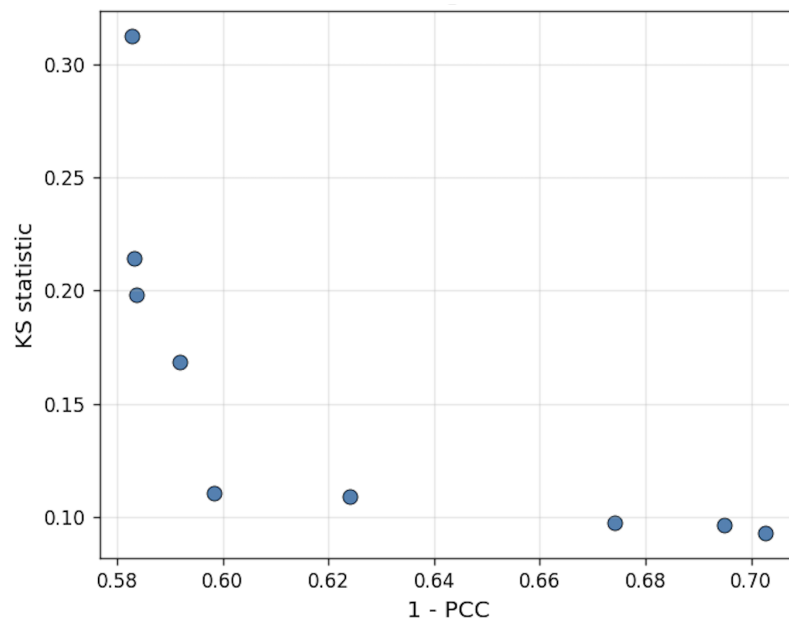


Figura 4.17: Fronte di Pareto ottenuto per un soggetto scelto randomicamente dalla coorte nell'ottimizzazione basata su PCC e KS. Ogni punto rappresenta una soluzione non dominata individuata dall'algoritmo genetico nello spazio degli obiettivi 1-PCC e KS.

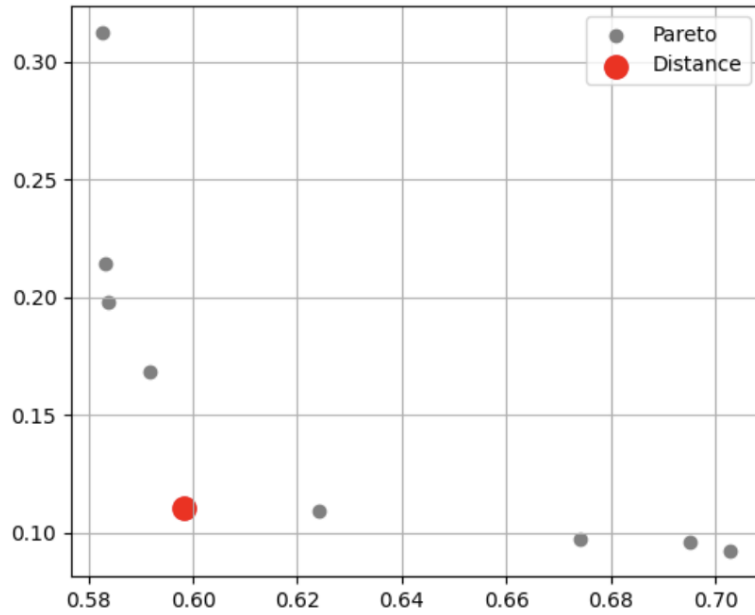


Figura 4.18: Selezione del punto ottimo sul fronte di Pareto del soggetto scelto randomicamente in 4.17 sopra. Il punto evidenziato rappresenta la soluzione scelta come miglior compromesso tra gli obiettivi 1-PCC e KS, secondo il criterio della distanza dal punto ideale.

Osservando il fronte di Pareto ottenuto con l'ottimizzazione basata su PCC e KS, si nota che le soluzioni non sono distribuite in modo casuale, ma tendono a disporsi lungo il fronte evidenziando il compromesso tra i due obiettivi. Dopo 40 generazioni, il fronte di Pareto risulta costituito da 9 soluzioni, che appaiono ben distribuite lungo il fronte, indicando che l'algoritmo è riuscito a preservare una buona diversità. La distribuzione delle soluzioni lungo il fronte suggerisce che entrambi gli obiettivi contribuiscano in modo comparabile alla selezione delle soluzioni. Infatti, non si osserva una prevalenza evidente di $1 - PCC$ o di KS , ma un compromesso tra le due metriche.

Ottimizzazione feature-based

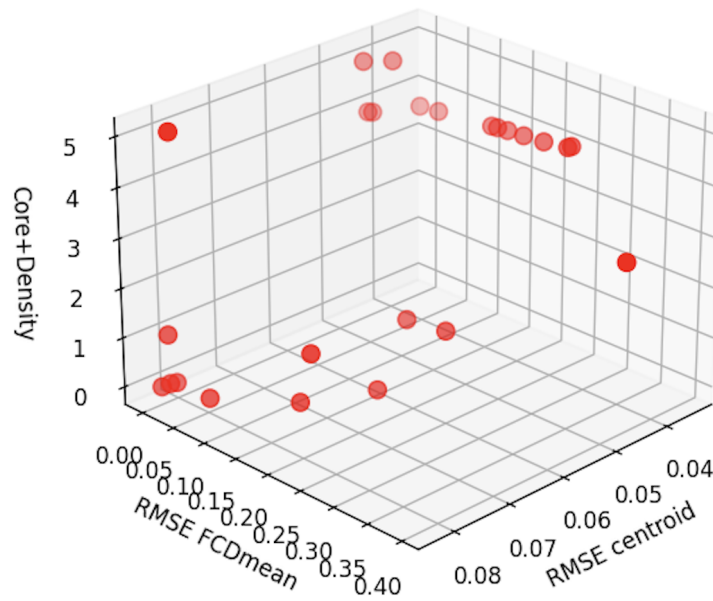


Figura 4.19: Fronte di Pareto finale tridimensionale ottenuto per un soggetto scelto randomicamente dalla coorte mediante ottimizzazione feature-based. Ogni punto del grafico rappresenta una soluzione nello spazio degli obiettivi definiti dalle feature considerate, ovvero RMSE del centroide, RMSE del valor medio di FCD e RMSE complessivo di nodi core e densità.

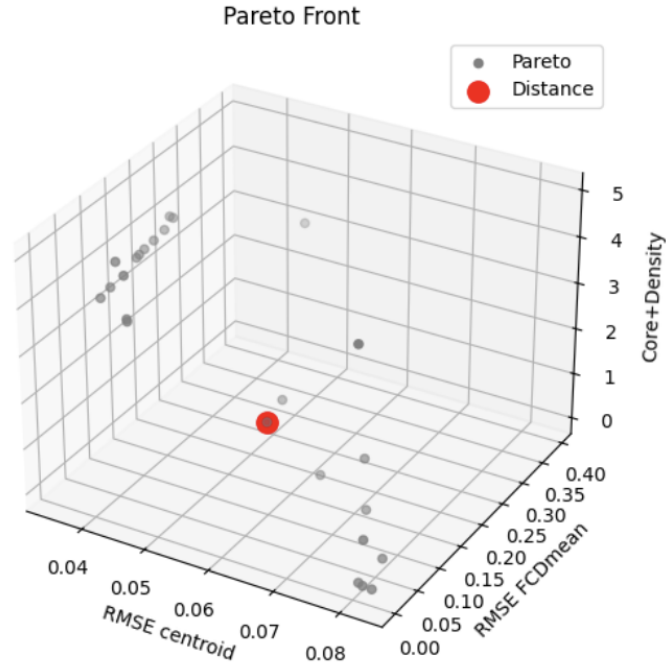


Figura 4.20: Selezione del punto ottimo sul fronte di Pareto feature-based del soggetto scelto randomicamente in 4.19. Il punto scelto rappresenta il miglior compromesso tra i tre obiettivi feature-based, secondo il criterio della distanza dal punto ideale.

Nel corso delle generazioni il fronte tende progressivamente a spostarsi verso regioni caratterizzate da valori più favorevoli degli obiettivi da minimizzare. Allo stesso tempo il numero di soluzioni può aumentare, descrivendo in maniera più completa i diversi compromessi tra gli obiettivi considerati.

Il fronte tridimensionale ottenuto con l'approccio feature-based mostra una struttura più complessa rispetto al fronte bidimensionale. In questo caso le soluzioni devono rappresentare un compromesso tra tre obiettivi differenti. Rispetto al fronte bidimensionale, il fronte tridimensionale permette quindi di descrivere la dinamica simulata in modo più ampio, perché non si limita al confronto tra matrici di connettività, ma include anche proprietà più specifiche del segnale BOLD e dell'organizzazione funzionale della rete. Tuttavia, proprio per la presenza di un obiettivo aggiuntivo, la distribuzione delle soluzioni risulta meno immediata da interpretare visivamente e il compromesso tra gli obiettivi appare più complesso.

4.2.4 Confronto tra gruppi clinici

Dopo aver trovato il punto ottimo per ciascun soggetto, i valori dei parametri sono stati raggruppati per gruppo clinico per valutare se ci sono differenze statisticamente significative tra i gruppi.

Ottimizzazione NSGA-II basata su PCC e KS

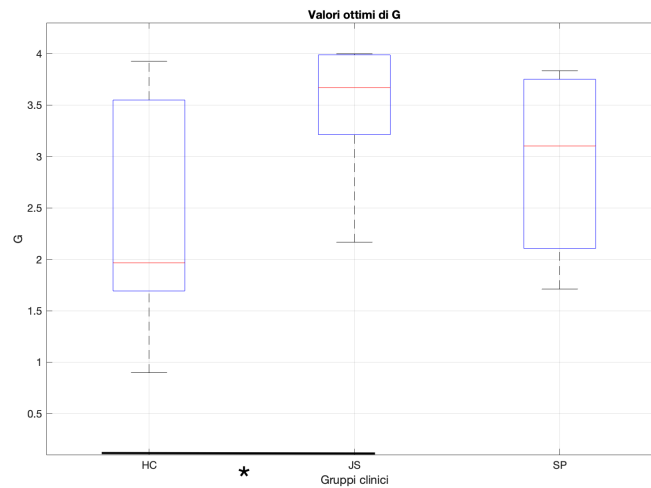


Figura 4.21: Boxplot dei parametri ottimi G suddivisi per gruppo clinico. La linea rossa interna indica la mediana, mentre i limiti inferiore e superiore del box rappresentano rispettivamente il primo e il terzo quartile. I baffi mostrano l'estensione dei valori osservati, mentre eventuali punti esterni indicano la presenza di outlier. * indica che c'è una differenza statisticamente significativa tra i gruppi considerati.

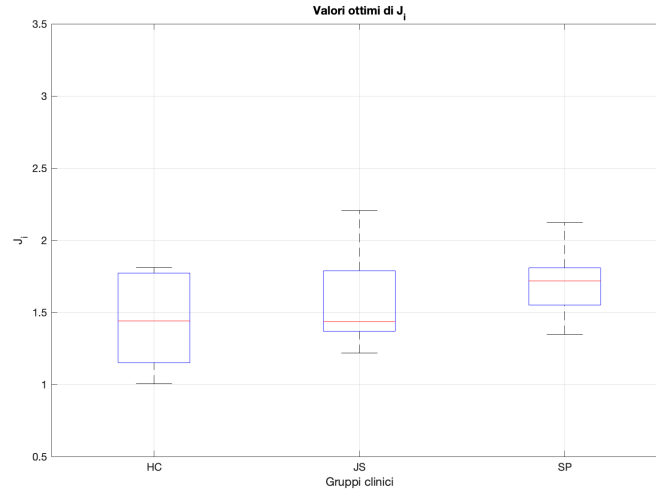


Figura 4.22: Boxplot dei parametri ottimi J_i suddivisi per gruppo clinico. La linea rossa interna indica la mediana, mentre i limiti inferiore e superiore del box rappresentano rispettivamente il primo e il terzo quartile. I baffi mostrano l'estensione dei valori osservati, mentre eventuali punti esterni indicano la presenza di outlier.

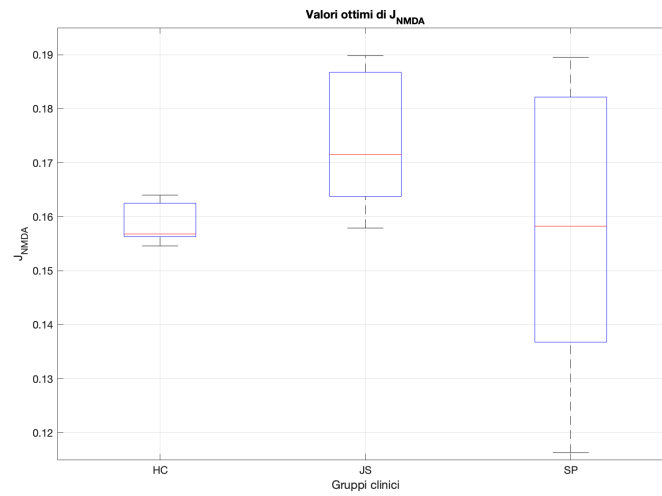


Figura 4.23: Boxplot dei parametri ottimi J_N suddivisi per gruppo clinico. La linea rossa interna indica la mediana, mentre i limiti inferiore e superiore del box rappresentano rispettivamente il primo e il terzo quartile. I baffi mostrano l'estensione dei valori osservati, mentre eventuali punti esterni indicano la presenza di outlier.

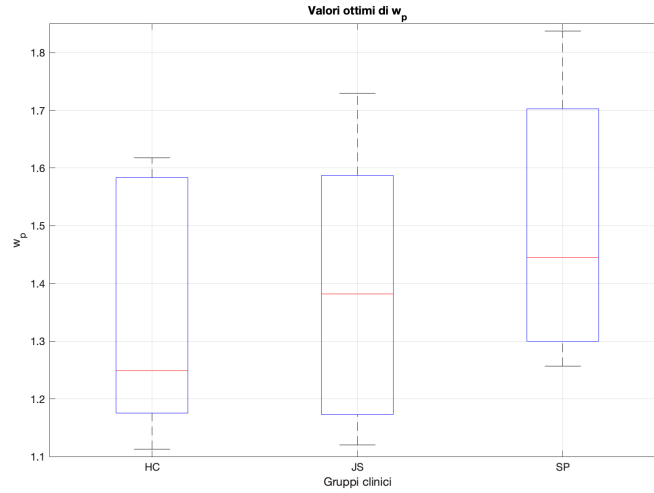


Figura 4.24: Boxplot dei parametri ottimi w_p suddivisi per gruppo clinico. La linea rossa interna indica la mediana, mentre i limiti inferiore e superiore del box rappresentano rispettivamente il primo e il terzo quartile. I baffi mostrano l'estensione dei valori osservati, mentre eventuali punti esterni indicano la presenza di outlier.

Ottimizzazione NSGA-II feature-based

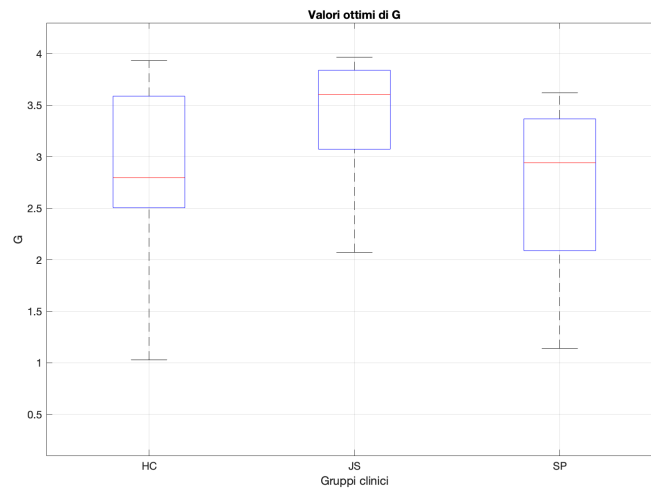


Figura 4.25: Boxplot dei parametri ottimi G suddivisi per gruppo clinico. La linea rossa interna indica la mediana, mentre i limiti inferiore e superiore del box rappresentano rispettivamente il primo e il terzo quartile. I baffi mostrano l'estensione dei valori osservati, mentre eventuali punti esterni indicano la presenza di outlier.

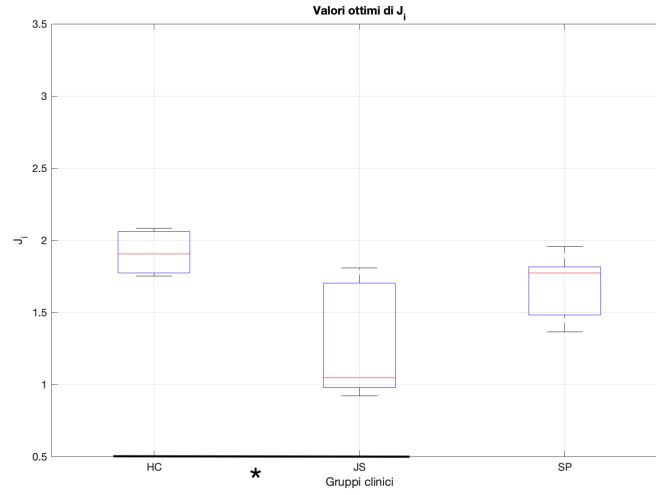


Figura 4.26: Boxplot dei parametri ottimi J_i suddivisi per gruppo clinico. La linea rossa interna indica la mediana, mentre i limiti inferiore e superiore del box rappresentano rispettivamente il primo e il terzo quartile. I baffi mostrano l'estensione dei valori osservati, mentre eventuali punti esterni indicano la presenza di outlier. * indica che c'è una differenza statisticamente significativa tra i gruppi considerati.

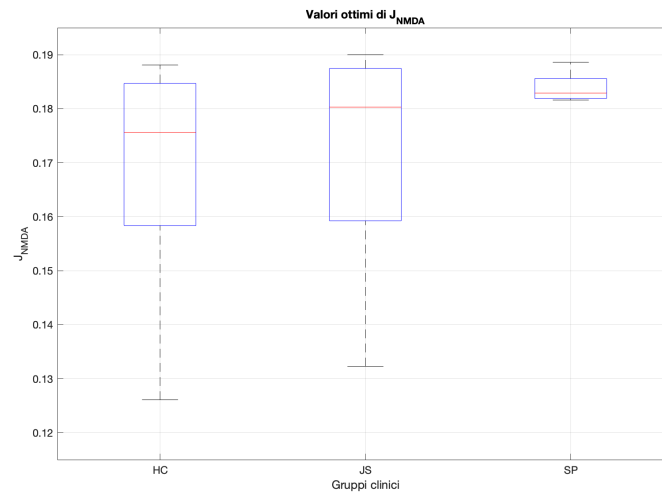


Figura 4.27: Boxplot dei parametri ottimi J_N suddivisi per gruppo clinico. La linea rossa interna indica la mediana, mentre i limiti inferiore e superiore del box rappresentano rispettivamente il primo e il terzo quartile. I baffi mostrano l'estensione dei valori osservati, mentre eventuali punti esterni indicano la presenza di outlier.

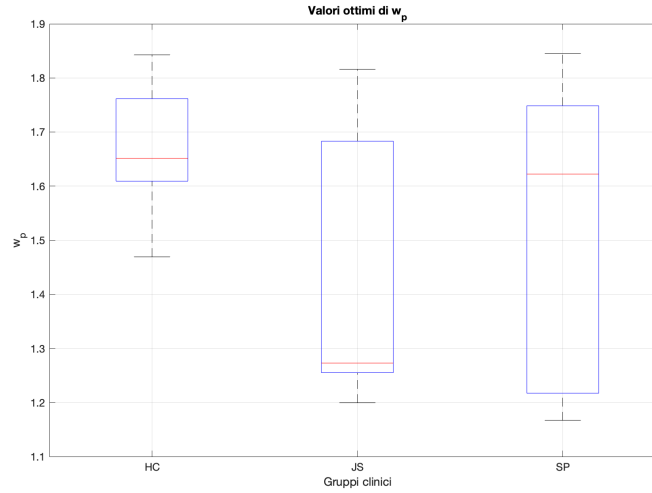


Figura 4.28: Boxplot dei parametri ottimi w_p suddivisi per gruppo clinico. La linea rossa interna indica la mediana, mentre i limiti inferiore e superiore del box rappresentano rispettivamente il primo e il terzo quartile. I baffi mostrano l'estensione dei valori osservati, mentre eventuali punti esterni indicano la presenza di outlier.

Successivamente, è stato applicato il test non parametrico di Kruskal-Wallis per il confronto globale tra i tre gruppi. Nei casi in cui il test ha evidenziato una differenza significativa, sono stati eseguiti confronti post-hoc mediante test di Dunn con correzione di Bonferroni.

Oltre ai risultati ottenuti con l'algoritmo genetico, sono stati considerati anche i risultati derivanti dall'ottimizzazione mediante grid-search. Tale approccio è stato utilizzato come riferimento, poiché è basato sugli stessi obiettivi (1-PCC e KS) ed è stato applicato allo stesso dataset. In questo caso il test di Kruskal-Wallis ha evidenziato una differenza statisticamente significativa tra i gruppi per il parametro G . Il test post-hoc di Dunn, corretto con metodo Bonferroni, ha indicato una differenza significativa tra i controlli sani e i pazienti con sindrome di Joubert.

Per l'ottimizzazione basata su PCC e KS, il test di Kruskal-Wallis ha evidenziato anche in questo caso una differenza statisticamente significativa tra i gruppi per il parametro G . Il test post-hoc di Dunn, corretto con metodo Bonferroni, ha indicato una differenza significativa tra i gruppi HC e JS. Per gli altri parametri non sono state osservate differenze statisticamente significative tra i gruppi.

Per quanto riguarda l'ottimizzazione feature-based, il test di Kruskal-Wallis

ha mostrato una differenza statisticamente significativa tra i gruppi per il parametro J_i . Il confronto post-hoc mediante test di Dunn con correzione di Bonferroni ha evidenziato una differenza significativa sempre tra i gruppi JS e HC. Per i restanti parametri dell'ottimizzazione feature-based non sono emerse differenze significative tra i gruppi. Per il gruppo clinico costituito dai soggetti con atassia lentamente progressiva non emergono differenze significative rispetto agli altri gruppi in entrambi i metodi di ottimizzazione. Questo indica che i valori ottenuti per questo gruppo si distribuiscono in modo sovrapposto a quelli degli altri gruppi, senza evidenziare una separazione statisticamente significativa.

Metodo	Parametro	H	p-value	Mediana HC	JS	SP
Grid-search PCC-KS	G	6.8951	0.0318	2.4000	3.7000	3.4500
Grid-search PCC-KS	J_i	1.0898	0.5799	1.5000	1.5500	1.5500
Grid-search PCC-KS	J_{NMDA}	5.0661	0.0794	0.1400	0.1525	0.1625
Grid-search PCC-KS	w_p	4.1401	0.1262	1.5000	1.3750	1.2750
NSGA-II PCC-KS	G	6.0114	0.0495	1.9653	3.6695	3.1017
NSGA-II PCC-KS	J_i	2.2422	0.3259	1.4403	1.4355	1.7180
NSGA-II PCC-KS	J_{NMDA}	3.1795	0.2040	0.1568	0.1715	0.1583
NSGA-II PCC-KS	w_p	2.6023	0.2722	1.2490	1.3818	1.4449
NSGA-II Feature-based	G	3.2956	0.1925	2.7985	3.6025	2.9417
NSGA-II Feature-based	J_i	9.7841	0.0075	1.9049	1.0468	1.6831
NSGA-II Feature-based	J_{NMDA}	2.2557	0.3237	0.1756	0.1803	0.1825
NSGA-II Feature-based	w_p	2.5633	0.2776	1.6511	1.2728	1.6223

Tabella 4.1: Risultati del test di Kruskal-Wallis applicato ai punti ottimi.

Metodo	Parametro	Confronto	p Bonferroni
Grid-search PCC-KS	G	JS vs HC	0.0277
NSGA-II PCC-KS	G	JS vs HC	0.0429
NSGA-II Feature-based	J_i	JS vs HC	0.0082

Tabella 4.2: Confronti post-hoc significativi ottenuti mediante test di Dunn con correzione di Bonferroni.

Capitolo 5

Discussione

In questo studio è stata effettuata l'ottimizzazione multi-obiettivo dei parametri biofisici del modello di Wong-Wang implementato in The Virtual Brain, con lo scopo di ottenere delle dinamiche simulate il più possibile confrontabili con i dati empirici. Per l'ottimizzazione dei parametri sono stati utilizzati due approcci differenti: il primo è un approccio standard basato su una grid search che ottimizza i parametri a coppie, il secondo invece utilizza un algoritmo genetico in grado di ottimizzare contemporaneamente tutti e quattro i parametri del modello. Per quest'ultimo approccio sono state considerate due diverse strategie di ottimizzazione: la prima prevede il confronto tra le matrici di connettività funzionale, sia statica che dinamica, per minimizzare i due obiettivi, che sono rispettivamente 1-PCC e la distanza KS, mentre la seconda utilizza misure quantitative, utilizzate per descrivere proprietà specifiche della dinamica cerebrale e della connettività funzionale.

I risultati ottenuti mostrano che la scelta delle funzioni obiettivo influenza in modo rilevante la distribuzione dei parametri ottimizzati e di conseguenza anche l'interpretazione della dinamica simulata. Questo avviene perché la funzione obiettivo definisce quali aspetti della dinamica devono essere riprodotti con maggiore accuratezza dal modello. Quindi, i parametri ottimizzati devono essere interpretati come soluzioni dipendenti dagli obiettivi scelti per l'ottimizzazione.

A differenza dell'approccio standard basato sulla grid search, che restituisce una singola combinazione del valore ottimale dei parametri, l'algoritmo genetico NSGA-II non restituisce un solo punto, ma un insieme di soluzioni non dominate che vanno a costituire il fronte di Pareto. Ciascuna soluzione

del fronte rappresenta un compromesso tra gli obiettivi, poichè un punto può riprodurre in modo migliore un obiettivo ma essere meno efficace in un altro. Per questo motivo non esiste necessariamente una soluzione migliore per tutti gli obiettivi.

Questo aspetto evidenzia proprio la natura multi-obiettivo del problema di ottimizzazione. Per esempio, il segnale BOLD simulato deve riprodurre contemporaneamente differenti aspetti dei dati empirici, tra cui la connettività funzionale statica, la connettività funzionale dinamica e specifiche feature riassuntive del segnale.

Il fronte di Pareto non rappresenta quindi una limitazione dell'algoritmo, poichè la presenza di più soluzioni non implica che una sola sia corretta e che le altre siano errate, ma mostra l'esistenza di diverse combinazioni di parametri in grado di descrivere aspetti complementari della dinamica cerebrale.

Analizzando la distribuzione dei parametri ottimizzati, emerge in generale una variabilità tra i soggetti. Questo risultato è atteso, poiché le simulazioni in TVB sono soggetto-specifiche: per ciascun soggetto il modello utilizza come input la propria matrice di connettività strutturale, ricostruita a partire dai dati di diffusione. Di conseguenza, la variabilità dei parametri non deve essere interpretata come un indice di instabilità dell'algoritmo, ma come il risultato dell'adattamento del modello alla specifica organizzazione strutturale e alla dinamica empirica di ciascun soggetto.

Tra i parametri analizzati assume un ruolo importante il parametro G dell'accoppiamento globale, poiché regola il contributo di connettività strutturale nella dinamica simulata, come descritto nell'equazione di Wong-Wang relativa alle correnti sinaptiche ricevute dalla popolazione eccitatoria di ciascun nodo. Al variare di questo parametro il modello necessita di un livello diverso di accoppiamento tra regioni cerebrali per riprodurre la dinamica funzionale. Gli altri tre parametri, ovvero J_i , J_N e w_p , contribuiscono a modulare il bilanciamento locale tra dinamiche inibitorie ed eccitatorie.

Inoltre, la variabilità dei parametri si può notare non solo tra soggetti diversi, ma anche confrontando gli stessi soggetti nei due metodi di ottimizzazione. Questo può essere spiegato dal fatto che le due strategie ottimizzano obiettivi differenti e quindi guidano l'algoritmo verso diverse combinazioni di parametri.

Per rendere possibile un confronto più diretto tra soggetti, gruppi clinici e metodi, è stato quindi necessario selezionare, per ciascun fronte di Pareto, una singola soluzione rappresentativa. Per farlo si è deciso di adottare come criterio

di selezione la distanza dal punto ideale. Questo approccio ha il vantaggio di poter essere applicato con la stessa logica a fronti con un diverso numero di obiettivi. In questo modo la soluzione migliore corrisponde al punto del fronte più vicino alla minimizzazione simultanea degli obiettivi considerati.

Una volta selezionati i punti ottimi per ciascun soggetto si sono confrontate le distribuzioni dei parametri tra i gruppi clinici. Dai risultati è emerso che nei metodi di ottimizzazione finalizzati alla minimizzazione degli obiettivi 1-PCC e KS, utilizzando sia l'algoritmo genetico che la grid search, il parametro G è l'unico che mostra una differenza statisticamente significativa tra i gruppi, in particolare tra i controlli sani e i soggetti con la sindrome di Joubert. G mostra valori più elevati nel gruppo JS, indicando che il modello necessita di un accoppiamento globale maggiore per riprodurre dinamiche cerebrali simili a quelle osservate empiricamente. Inoltre questa differenza potrebbe essere anche legata ad alterazioni della connettività strutturale associate alla patologia. Poiché in TVB la connettività strutturale definisce le connessioni tra i nodi della rete, un aumento di G potrebbe rappresentare un adattamento del modello necessario a compensare una diversa organizzazione delle connessioni strutturali.

Per quanto riguarda l'ottimizzazione basata su feature, il parametro che in questo metodo mostra una differenza statisticamente significativa tra i gruppi è il J_i , associato all'accoppiamento inibitorio locale, sempre tra i controlli sani e i soggetti con la sindrome di Joubert. In questo caso, i valori più elevati osservabili nei controlli sani suggeriscono che il modello richieda un maggiore contributo dell'inibizione locale rispetto al gruppo JS per minimizzare la differenza tra feature empiriche e simulate.

Per i soggetti con atassia lentamente progressiva, in entrambi i metodi di ottimizzazione, non emergono differenze statisticamente significative rispetto agli altri gruppi. Questo risultato può essere interpretato considerando che, in questo gruppo, l'alterazione della dinamica cerebrale potrebbe non riflettersi principalmente in modifiche dei parametri di accoppiamento globale o di accoppiamento inibitorio locale. La mancata significatività potrebbe dipendere sia da una maggiore variabilità inter-soggetto sia dalla dimensione ridotta della coorte, che limita la possibilità di rilevare differenze meno marcate tra gruppi.

L'approccio basato su feature permette di aggiungere informazioni diverse rispetto al solo confronto tra matrici di connettività. Infatti, le feature selezionate descrivono proprietà specifiche segnale BOLD e della rete funzionale.

Questo rappresenta un vantaggio perché tali indici entrano direttamente nella funzione obiettivo e contribuiscono quindi a guidare l'ottimizzazione. Al contrario, né la grid search né l'ottimizzazione basata su PCC e KS utilizzano direttamente informazioni sul segnale BOLD, ma solo misure derivate dalla connettività. In questo modo, le feature permettono di spostare l'ottimizzazione da un semplice confronto tra matrici verso una descrizione più mirata della dinamica simulata del modello.

Infine si è passati ad analizzare le diverse strategie di ottimizzazione. Il confronto tra la grid search e l'algoritmo genetico NSGA-II ha mostrato differenze statisticamente significative nella minimizzazione della funzione obiettivo 1-PCC. Quindi questo risultato indica che NSGA-II riesce a ottenere una maggiore somiglianza tra la connettività funzionale statica simulata e quella empirica. Questo risultato era in parte atteso, considerando che la grid search presenta il limite di ottimizzare i parametri a due a due, e non simultaneamente tutti e quattro come l'algoritmo genetico. In generale NSGA-II ha il vantaggio di esplorare lo spazio dei parametri in modo più flessibile e di ottimizzare contemporaneamente l'intera combinazione dei parametri.

Un altro confronto è stato fatto tra il metodo basato su PCC e KS e il metodo basato sulle feature. Come previsto i due approcci portano a risultati differenti perché non ottimizzano le stesse grandezze, tendendo a selezionare combinazioni di parametri agli obiettivi considerati. Questa differenza conferma che la scelta della funzione obiettivo influenza il risultato finale dell'ottimizzazione.

Inoltre questo modo diverso di valutare la dinamica simulata può aiutare anche a interpretare i parametri risultati significativi nei due metodi. Nel caso dell'ottimizzazione basata su PCC e KS, la significatività di G è coerente con il fatto che questi obiettivi valutano la somiglianza tra matrici di connettività funzionale, quindi aspetti legati all'interazione tra le regioni della rete. Nel metodo feature-based sono significative le differenze per il parametro J_i , che regola l'accoppiamento inibitorio locale. Questo può suggerire che le feature selezionate siano più sensibili a proprietà locali del modello e al bilanciamento tra attività eccitatoria e inibitoria.

Per confrontare meglio i due approcci, sono stati poi calcolati a posteriori i valori di 1-PCC e KS anche per i punti ottimi ottenuti con il metodo feature-based, e viceversa sono state calcolate le feature anche per i punti ottimi ottenuti con il metodo basato su PCC e KS. In questo modo è stato possibile

valutare se un metodo riuscisse comunque a mantenere buone prestazioni anche rispetto a metriche non usate direttamente durante l'ottimizzazione. Il confronto mostra che i due approcci non sono equivalenti.

Come prospettiva futura, questo lavoro potrebbe essere esteso considerando reti cerebrali differenti, in modo da valutare se le differenze osservate siano specifiche della rete attentiva ventrale oppure presenti anche in altre reti funzionali. Quindi potrebbe essere utile analizzare reti maggiormente coinvolte nello studio delle atassie, come reti motorie, somatomotorie e fronto-parietali. Un ulteriore sviluppo potrebbe riguardare l'aggiunta di nuove feature, per verificare se l'inclusione di ulteriori misure possa contribuire a descrivere in modo più completo la dinamica cerebrale simulata. Un altro aspetto da migliorare potrebbe essere quello di includere un numero maggiore di partecipanti. È importante sottolineare che, nel contesto di TVB, l'ottimizzazione è soggetto-specifica e quindi la simulazione del singolo soggetto non dipende direttamente dalla dimensione del campione, ma un numero maggiore di soggetti permetterebbe di ottenere confronti statistici più robusti tra i gruppi clinici e di valutare con maggiore affidabilità la presenza di differenze nei parametri ottimizzati.

Nel complesso, i risultati ottenuti mostrano che NSGA-II può essere uno strumento utile per ottimizzare i parametri di The Virtual Brain e per confrontare diverse strategie di fitting della dinamica cerebrale.

Capitolo 6

Conclusioni

Questa tesi si è concentrata sull'ottimizzazione dei parametri biofisici di un modello implementato in The Virtual Brain, con lo scopo di minimizzare specifiche funzioni obiettivo e quindi di ottenere una dinamica cerebrale simulata il più possibile confrontabile con i dati empirici.

A questo scopo è stato utilizzato NSGA-II, un algoritmo genetico che consente di esplorare simultaneamente lo spazio di tutti e quattro i parametri del modello e ha consentito di individuare per ciascun soggetto un insieme di soluzioni non dominate costituenti il fronte di Pareto. Sono state utilizzate due strategie di ottimizzazione che hanno portato il modello a selezionare combinazioni di parametri differenti.

Un ulteriore confronto è stato effettuato con il metodo standard basato su grid search, in cui i parametri vengono ottimizzati a due a due. Questo confronto è stato utile per verificare se NSGA-II fosse effettivamente più efficace nell'individuare combinazioni parametriche in grado di minimizzare gli obiettivi considerati. Dai risultati è emerso che l'approccio basato su NSGA-II permette di ottenere prestazioni migliori rispetto alla grid search, in particolare per la funzione obiettivo 1-PCC. Questo indica che la connettività funzionale statica simulata risulta più simile a quella empirica quando i parametri vengono ottimizzati tramite l'algoritmo genetico.

In conclusione, questa studio ha permesso di sviluppare una procedura di ottimizzazione multiparametrica per simulazioni BOLD implementate in TVB con il modello di Wong-Wang. L'approccio proposto migliora il metodo standard basato su grid search, poiché consente di esplorare simultaneamente tutti i parametri. Pur essendo stato applicato in questo studio a una coorte

comprendente soggetti con alterazioni neurologiche, il metodo non è specifico per una singola patologia, ma rappresenta uno strumento flessibile che può essere esteso ad altri studi basati su simulazioni BOLD.

Bibliografia

- [1] Donald B. Plewes and Walter Kucharczyk. Physics of mri: A primer. *Journal of Magnetic Resonance Imaging*, 35(5):1038–1054, 2012. doi: 10.1002/jmri.23642.
- [2] Bernd André Jung and Matthias Weigel. Spin echo magnetic resonance imaging. *Journal of Magnetic Resonance Imaging*, 37(4):805–817, 2013. doi: 10.1002/jmri.24068.
- [3] Seiji Ogawa, Tso-Ming Lee, Alan R. Kay, and David W. Tank. Brain magnetic resonance imaging with contrast dependent on blood oxygenation. *Proceedings of the National Academy of Sciences of the United States of America*, 87(24):9868–9872, 1990. doi: 10.1073/pnas.87.24.9868.
- [4] Kenneth K. Kwong, John W. Belliveau, David A. Chesler, Iris E. Goldberg, Robert M. Weisskoff, Bernard P. Poncelet, David N. Kennedy, Bruce E. Hoppel, Mark S. Cohen, Robert Turner, Hsing-Mao Cheng, Thomas J. Brady, and Bruce R. Rosen. Dynamic magnetic resonance imaging of human brain activity during primary sensory stimulation. *Proceedings of the National Academy of Sciences of the United States of America*, 89(12):5675–5679, 1992. doi: 10.1073/pnas.89.12.5675.
- [5] Peter A. Bandettini. Twenty years of functional mri: The science and the stories. *NeuroImage*, 62(2):575–588, 2012. doi: 10.1016/j.neuroimage.2012.04.026.
- [6] Gary H. Glover. Overview of functional magnetic resonance imaging. *Neurosurgery Clinics of North America*, 22(2):133–139, 2011. doi: 10.1016/j.nec.2010.11.001.

- [7] Nikos K. Logothetis, Jon Pauls, Mark Augath, Torsten Trinath, and Axel Oeltermann. Neurophysiological investigation of the basis of the fmri signal. *Nature*, 412(6843):150–157, 2001. doi: 10.1038/35084005.
- [8] Gary H. Glover. Overview of functional magnetic resonance imaging. *Neurosurgery Clinics of North America*, 22(2):133–139, 2011. doi: 10.1016/j.nec.2010.11.001.
- [9] Bharat Biswal, F. Zerrin Yetkin, Victor M. Haughton, and James S. Hyde. Functional connectivity in the motor cortex of resting human brain using echo-planar mri. *Magnetic Resonance in Medicine*, 34(4):537–541, 1995. doi: 10.1002/mrm.1910340409.
- [10] Paula Sanz Leon, Stuart A. Knock, Marmaduke M. Woodman, Lia Domide, Jochen Mersmann, Anthony R. McIntosh, and Viktor K. Jirsa. The virtual brain: a simulator of primate brain network dynamics. *Frontiers in Neuroinformatics*, 7:10, 2013. doi: 10.3389/fninf.2013.00010.
- [11] Paula Sanz-Leon, Stuart A. Knock, Andreas Spiegler, and Viktor K. Jirsa. Mathematical framework for large-scale brain network modeling in the virtual brain. *NeuroImage*, 111:385–430, 2015. doi: 10.1016/j.neuroimage.2015.01.002.
- [12] Kong-Fatt Wong and Xiao-Jing Wang. A recurrent network mechanism of time integration in perceptual decisions. *Journal of Neuroscience*, 26(4):1314–1328, 2006. doi: 10.1523/JNEUROSCI.3733-05.2006.
- [13] Gustavo Deco, Adrian Ponce-Alvarez, Patric Hagmann, Gian Luca Romani, Dante Mantini, and Maurizio Corbetta. How local excitation-inhibition ratio impacts the whole brain dynamics. *Journal of Neuroscience*, 34(23):7886–7898, 2014. doi: 10.1523/JNEUROSCI.5068-13.2014.
- [14] Anita Monteverdi, Fulvia Palesi, Alfredo Costa, Paolo Vitali, Anna Pichiecchio, Matteo Cotta Ramusino, Sara Bernini, Viktor Jirsa, Claudia A. M. Gandini Wheeler-Kingshott, and Egidio D’Angelo. Subject-specific features of excitation/inhibition profiles in neurodegenerative diseases. *Frontiers in Aging Neuroscience*, 14:868342, 2022. doi: 10.3389/fnagi.2022.868342.

- [15] Kalyanmoy Deb, Amrit Pratap, Sameer Agarwal, and T. Meyarivan. A fast and elitist multiobjective genetic algorithm: Nsga-ii. *IEEE Transactions on Evolutionary Computation*, 6(2):182–197, 2002. doi: 10.1109/4235.996017.
- [16] Marta Gaviraghi, Anita Monteverdi, Sara Bulgheroni, Marta Mercati, Arianna De Laurentiis, Anna Nigri, Marina Grisoli, Stefano D’Arrigo, Claudia A. M. Gandini Wheeler-Kingshott, Claudia Casellato, Fulvia Palesi, and Egidio D’Angelo. Brain digital twins reveal network changes in congenital and slowly progressive cerebellar ataxias. *bioRxiv*, 2026. doi: 10.64898/2026.03.23.713380. Preprint.
- [17] Maurizio Corbetta and Gordon L. Shulman. Control of goal-directed and stimulus-driven attention in the brain. *Nature Reviews Neuroscience*, 3(3):201–215, 2002. doi: 10.1038/nrn755.
- [18] Maurizio Corbetta, Gaurav Patel, and Gordon L. Shulman. The reorienting system of the human brain: From environment to theory of mind. *Neuron*, 58(3):306–324, 2008. doi: 10.1016/j.neuron.2008.04.017.
- [19] Mark Jenkinson, Christian F. Beckmann, Timothy E. J. Behrens, Mark W. Woolrich, and Stephen M. Smith. FSL. *NeuroImage*, 62(2):782–790, 2012. doi: 10.1016/j.neuroimage.2011.09.015.
- [20] Anita Monteverdi, Fulvia Palesi, Michael Schirner, Francesca Argentino, Mariateresa Merante, Alberto Redolfi, Francesca Conca, Laura Mazzocchi, Stefano F. Cappa, Matteo Cotta Ramusino, Alfredo Costa, Anna Pichiecchio, Lisa M. Farina, Viktor Jirsa, Petra Ritter, Claudia A. M. Gandini Wheeler-Kingshott, and Egidio D’Angelo. Virtual brain simulations reveal network-specific parameters in neurodegenerative dementias. *Frontiers in Aging Neuroscience*, 15:1204134, 2023. doi: 10.3389/fnagi.2023.1204134.
- [21] Julian Blank and Kalyanmoy Deb. pymoo: Multi-objective optimization in python. *IEEE Access*, 8:89497–89509, 2020.