



ALMA MATER STUDIORUM · UNIVERSITÀ DI BOLOGNA

Dipartimento di Informatica — Scienza e Ingegneria
Corso di Laurea Triennale in Informatica

Deep Image Prior Vincolato per la Ricostruzione Tomografica a Proiezioni Limitate

Relatore:
Chiar.ma Prof.ssa
Elena Loli Piccolomini

Presentata da:
Leonardo Verrone

VI Sessione
Anno Accademico 2024/2025

Sommario

La Tomografia Computerizzata a raggi X produce misurazioni indirette della struttura interna dell'oggetto esaminato, e la ricostruzione dell'immagine a partire da questi dati costituisce un problema inverso mal posto, la cui risoluzione richiede tecniche di regolarizzazione adeguate. I metodi supervisionati basati sul Deep Learning hanno mostrato risultati eccellenti in questo contesto, ma dipendono dalla disponibilità di grandi dataset con immagini di riferimento affidabili, difficili da ottenere in ambito medico senza esporre il paziente a dosi di radiazioni inaccettabili.

In questo lavoro viene applicato alla ricostruzione tomografica 2D il metodo cDIP-RED, proposto da Cascarano et al., che si inserisce nel framework del Deep Image Prior (DIP). A differenza dei metodi supervisionati, l'approccio DIP non richiede alcun dataset di addestramento, poiché i pesi della rete neurale vengono ottimizzati direttamente sull'immagine da ricostruire. Il metodo cDIP-RED riformula il problema di ottimizzazione come un problema vincolato secondo il principio di discrepanza di Morozov, eliminando la necessità di determinare manualmente il parametro di regolarizzazione e rendendo il metodo robusto rispetto alla scelta del numero di iterazioni, criticità tipica del DIP classico basato sull'early stopping.

Il lavoro si basa su una implementazione del metodo cDIP-RED applicata alla ricostruzione di immagini tomografiche 2D in condizioni di dati limitati, di cui vengono discussi i risultati sperimentali a confronto con il metodo classico Filtered Back-Projection.

Introduzione

Al giorno d’oggi, molti problemi di interesse pratico richiedono di ricavare informazioni da misurazioni indirette o incomplete. L’avanzamento tecnologico degli ultimi decenni ha moltiplicato le sorgenti di dati disponibili, rendendo possibile misurare fenomeni un tempo inaccessibili. La criticità è che questi dati raramente descrivono direttamente il fenomeno che si vuole studiare, e per ricavarne una rappresentazione interpretabile è quindi necessario eseguire un processo di ricostruzione che risalga dagli effetti osservati alla loro causa. Questa tipologia di problemi prende il nome di problemi inversi, e la loro risoluzione è tutt’altro che banale in quanto risultano spesso mal posti, nel senso che piccole perturbazioni nei dati misurati possono produrre soluzioni molto diverse tra loro.

Negli ultimi anni, i metodi supervisionati basati sul *Deep Learning* (DL) hanno mostrato risultati eccezionali nella risoluzione di problemi inversi mal posti, grazie alla capacità delle reti neurali convoluzionali profonde (*ConvNets*) di apprendere rappresentazioni complesse direttamente dai dati [5]. Tuttavia, questi metodi richiedono una fase di addestramento su un ampio *dataset* composto da coppie formate da una misurazione e una rappresentazione dell’oggetto misurato (*ground truth*) [5], che non sempre sono disponibili o facili da ottenere in grandi quantità.

Questa limitazione è particolarmente critica nel campo dell’imaging medico. Nel caso della Tomografia Computerizzata a raggi X (TC) è molto difficile ottenere i dati *ground truth*, in quanto è impossibile acquisire direttamente un’immagine dell’interno del paziente. Per questo motivo i me-

todi DL solitamente considerano come ground truth un'immagine ottenuta utilizzando metodi classici di ricostruzione tomografica, come ad esempio il filtered back-projection (FBP), che però richiedono un'elevata quantità di misurazioni. Sarebbe quindi necessario sottoporre il paziente a dosi di radiazioni da raggi X inaccettabilmente elevate, rendendo estremamente difficile la creazione di dataset contenenti immagini di riferimento affidabili, rendendo problematica l'applicazione dei metodi di DL supervisionati [1].

In questi casi è opportuno utilizzare un'altra tipologia di metodi: quelli basati sull'approccio Deep Image Prior (DIP) [7], il quale, a differenza dell'approccio Deep Learning classico, non prevede l'utilizzo di un dataset per l'addestramento del modello, in quanto i pesi della rete neurale vengono ottimizzati direttamente su una singola immagine da ricostruire. I metodi DIP trovano pertanto applicazione negli ambiti in cui la disponibilità di dati per l'addestramento è limitata.

Durante questo lavoro di tesi applicherò l'approccio DIP alla ricostruzione di immagini tomografiche 2D, utilizzando il modello vincolato con regolarizzazione automatica presentato in [4]. Quest'ultimo riformula il problema di ottimizzazione introdotto dall'approccio DIP in un problema vincolato secondo il principio di discrepanza di Morozov, eliminando la necessità di determinare manualmente il valore del parametro di regolarizzazione.

La tesi è strutturata come segue. Nel Capitolo 1 viene illustrato il processo di Tomografia Computerizzata, descrivendo prima il meccanismo di acquisizione delle misurazioni tomografiche e il fenomeno fisico di attenuazione dei raggi X durante l'attraversamento di un corpo, per poi ricavarne la formulazione discreta che ne consente il trattamento matematico e computazionale. Nel Capitolo 2 viene descritto il metodo di ricostruzione basato sul machine learning oggetto di questa tesi, introducendo innanzitutto l'approccio Deep Image Prior e approfondendo in seguito il lavoro di Cascarano et al. [4] nello sviluppo di un modello di ottimizzazione vincolato con regolarizzazione automatica, di cui il metodo cDIP-RED viene illustrato in dettaglio nella Sezione 2.3. Infine, nel Capitolo 3 vengono riportati i risultati speri-

mentali ottenuti dall'implementazione sviluppata durante questo lavoro di tesi.

Indice

Sommario	i
Introduzione	iii
1 Tomografia Computerizzata	1
1.1 Modello matematico	2
1.1.1 Trasformata di Radon	3
1.1.2 Legge di Beer-Lambert	5
1.1.3 Discretizzazione e formulazione matriciale	6
2 Deep Image Prior e Modello di Ottimizzazione Vincolato	11
2.1 Deep Image Prior	12
2.2 Metodo PGDA	18
2.3 Modello cDIP-RED	20
2.4 Dettagli implementativi	27
3 Risultati sperimentali	33
3.1 Metriche di valutazione	33
3.2 Convergenza del modello	34
Conclusioni	39
Bibliografia	41

Capitolo 1

Tomografia Computerizzata

La Tomografia Computerizzata a raggi X (TC) è una delle tecnologie più preziose nell'imaging medico moderno [1], in quanto consente l'acquisizione non invasiva della struttura interna del corpo umano sfruttando la proprietà dei raggi X di attraversare la materia con diversi gradi di attenuazione [6]. Introdotta clinicamente negli anni '70, ha profondamente trasformato la diagnostica per immagini, e la ricerca attuale si concentra sempre di più sulla riduzione della dose di radiazioni a cui il paziente viene sottoposto durante la rilevazione [1].

La TC si basa su un esperimento concettualmente semplice. Una sorgente emette un fascio di raggi X che attraversa la sezione trasversale del paziente e viene rilevato da un detector lineare posto sul lato opposto. Il detector registra l'intensità residua dei raggi X dopo l'attraversamento dei tessuti, ottenendo così una proiezione dell'oggetto lungo quella direzione. Poiché una singola proiezione contiene informazioni troppo parziali sulla struttura interna, la coppia sorgente-detector ruota attorno al paziente e la misurazione viene ripetuta a diversi angoli, vedi Figura 1.1 [6, 9].

L'insieme di tutte le proiezioni così raccolte viene organizzato in una struttura bidimensionale chiamata *sinogramma*. Le proiezioni acquisite ai diversi angoli vengono accostate una accanto all'altra, ordinate per angolo crescente, formando un'unica immagine in cui ogni colonna corrisponde a

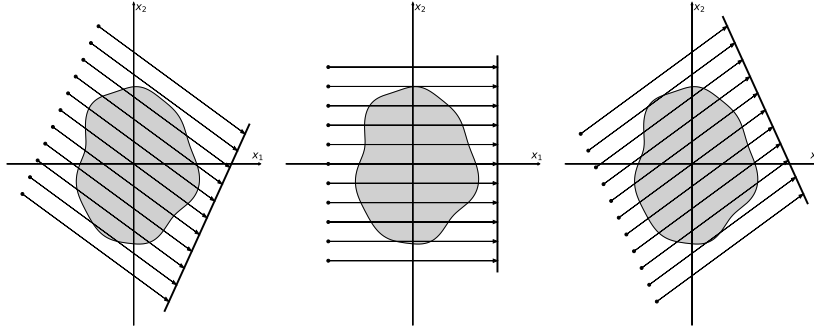


Figura 1.1: Esempio del processo di acquisizione tomografica. L'oggetto viene misurato più volte facendo ruotare la sorgente del fascio di raggi X e il rivelatore, quest'ultimo indicato con una linea spessa.

una proiezione [6]. Un esempio è mostrato in Figura 1.2, dove al phantom di Shepp-Logan viene affiancato il corrispondente sinogramma.

Il problema di ricostruzione consiste quindi nell'invertire questo processo di misura: dato il sinogramma, si cerca di determinare la distribuzione spaziale del coefficiente di attenuazione all'interno della sezione esaminata, che è la quantità fisica direttamente collegata alla composizione e alla densità dei tessuti attraversati [6].

1.1 Modello matematico

Il processo di acquisizione appena descritto richiede ora una formalizzazione matematica. In particolare, si vuole stabilire una relazione precisa tra le misure raccolte dal detector e la distribuzione del coefficiente di attenuazione all'interno dell'oggetto, ricavando la forma tipica di un problema lineare inverso:

$$Ax + \varepsilon = x_0,$$

dove x_0 è il vettore delle misurazioni raccolte dal detector (il sinogramma), x è il vettore incognito che rappresenta la distribuzione discreta del coefficiente di attenuazione, ε è il rumore di misurazione additivo e A è l'operatore di forward, che modella matematicamente il processo fisico di acquisizione delle

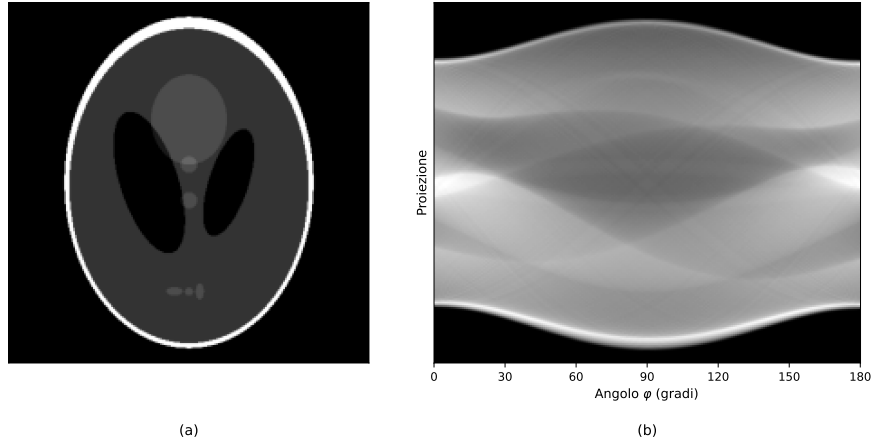


Figura 1.2: (a) Il phantom di Shepp-Logan, immagine di test che modella la sezione trasversale di un paziente. Il nero indica attenuazione nulla, mentre il bianco indica tessuti di attenuazione massima [6]. (b) il corrispondente sinogramma: sull'asse orizzontale gli angoli φ di proiezione, sull'asse verticale le proiezioni per ogni angolo.

proiezioni. A tal fine, verranno prima introdotti gli integrali di linea e la *Trasformata di Radon*, per poi ricavare la *Legge di Beer-Lambert* e, a partire da essa, costruire e discretizzare il modello di acquisizione, ottenendo la forma matriciale dell'operatore A .

1.1.1 Trasformata di Radon

È innanzitutto necessario introdurre una rappresentazione matematica sia dell'oggetto scansionato sia dei raggi X. Ci si riferisce al caso bidimensionale, in cui l'oggetto è descritto da una funzione positiva $f: \mathbb{R}^2 \rightarrow \mathbb{R}$ che rappresenta la distribuzione spaziale del coefficiente di attenuazione. I raggi X vengono invece modellizzati come rette nel piano, parametrizzate tramite la loro forma normale (φ, s) , dove $\varphi \in [0, \pi]$ è l'angolo del vettore normale alla retta, e $s \in \mathbb{R}$ è la distanza dal centro. L'equazione di un raggio (φ, s) è quindi:

$$x_1 \cos \varphi + x_2 \sin \varphi = s.$$

È utile inoltre descrivere il percorso del raggio (φ, s) al variare di $t \in \mathbb{R}$ come:

$$L_{\varphi,s}(t) = s\omega(\varphi) + t\omega^\perp(\varphi), \quad \omega(\varphi) := \begin{bmatrix} \cos \varphi \\ \sin \varphi \end{bmatrix}.$$

Nella Figura 1.3 viene mostrato il sistema di coordinate adottato, in cui due raggi (φ, s_1) e (φ, s_2) attraversano l'oggetto $f(x_1, x_2)$ e colpiscono il detector.

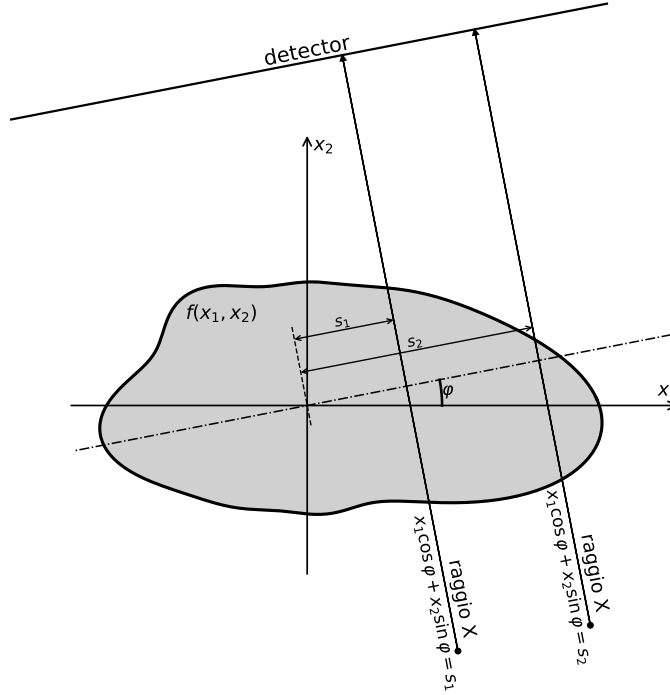


Figura 1.3: Geometria di riferimento nel piano (x_1, x_2) : due raggi X paralleli (φ, s_1) e (φ, s_2) che attraversano l'oggetto $f(x_1, x_2)$ e raggiungono il detector.

Un *integrale di linea* calcolato su un oggetto è l'integrale di alcune proprietà dell'oggetto stesso lungo una retta [6]. Dati un raggio X (φ, s) e un oggetto $f(x_1, x_2)$, l'integrale di linea è definito come:

$$P_\varphi(s) = \int_{\mathbb{R}} f(L_\varphi(t)) dt.$$

La funzione $P_\varphi(s)$ prende anche il nome di *Trasformata di Radon* di $f(x_1, x_2)$ [6]. Nel contesto della TC, $P_\varphi(s)$ quantifica esattamente la totale attenua-

zione subita dal raggio (φ, s) nell'attraversare l'oggetto: nella sezione successiva verrà mostrato come questa quantità si colleghi all'intensità rilevata dal detector.

Una *proiezione* è formata combinando un insieme di integrali di linea. La proiezione più semplice è quella ottenuta da un fascio di raggi X paralleli tra loro: scelto un angolo φ si calcola $P_\varphi(s)$ al variare di s . Ad esempio, nella Figura 1.4 viene mostrata la proiezione a raggi paralleli dell'oggetto $f(x_1, x_2)$ calcolata sull'angolo φ .

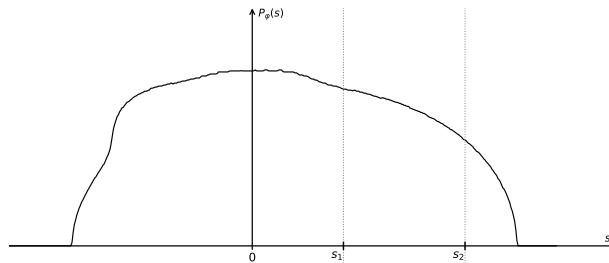


Figura 1.4: Proiezione a raggi paralleli dell'oggetto $f(x_1, x_2)$

1.1.2 Legge di Beer-Lambert

La trasformata di Radon da sola non è sufficiente a modellizzare il problema della ricostruzione tomografica, poiché è necessario stabilire una relazione precisa tra le grandezze fisiche misurabili e i coefficienti di attenuazione dell'oggetto. Durante una scansione tomografica i dati a disposizione sono essenzialmente due: l'intensità iniziale del fascio alla sorgente I_0 , nota tramite calibrazione, e l'intensità residua I_d rilevata dal detector dopo che il raggio ha attraversato il corpo.

Per descrivere matematicamente questo processo si considera un raggio X che attraversa uno spessore infinitesimo di materiale dx . In questo tratto il fascio subisce una riduzione della propria intensità dI proporzionale all'intensità corrente del raggio I e allo spessore attraversato, secondo una costante μ definita coefficiente di attenuazione lineare [6]. Tale fenomeno è espresso

dall'equazione differenziale

$$\frac{dI}{I} = -\mu \cdot dx, \quad (1.1)$$

la quale descrive la perdita relativa di intensità dovuta ai processi di assorbimento e diffusione.

Nell'ambito della TC bidimensionale, come descritto nella sezione precedente, l'oggetto è rappresentato da una funzione $f(x_1, x_2)$ che descrive la distribuzione spaziale del coefficiente di attenuazione. Poiché tale distribuzione varia da punto a punto lungo il cammino del raggio, il coefficiente μ non può essere trattato come una costante e viene quindi sostituito dalla funzione $f(x_1, x_2)$. La relazione differenziale (1.1) diventa quindi

$$\frac{dI}{I} = -f(x_1, x_2) dx.$$

Integrando entrambi i membri lungo l'intero cammino del raggio, dalla sorgente al detector, si ottiene

$$\int_{I_0}^{I_d} \frac{dI}{I} = - \int_{\mathbb{R}} f(L_{\varphi,s}(t)) dt,$$

da cui, risolvendo il membro sinistro, segue la *Legge di Beer-Lambert* [6]:

$$\ln \left(\frac{I_d}{I_0} \right) = - \int_{\mathbb{R}} f(L_{\varphi,s}(t)) dt$$

Questa relazione evidenzia come la distribuzione interna dell'oggetto $f(x_1, x_2)$ determini direttamente il rapporto di intensità misurato dal detector.

1.1.3 Discretizzazione e formulazione matriciale

Il passaggio dalla formulazione continua a una trattabile mediante algoritmi richiede una trasformazione dei parametri in termini discreti. Tale processo non costituisce soltanto una necessità computazionale, bensì riflette la realtà strumentale della scansione tomografica, in cui le posizioni della sorgente sono in numero finito. Risulta quindi indispensabile operare una

discretizzazione che riguardi sia i raggi X impiegati per la misurazione sia l'oggetto analizzato lungo la sezione trasversale.

La geometria di acquisizione viene inizialmente definita attraverso un insieme finito di $K > 0$ angoli di proiezione uniformemente distribuiti, indicati con φ_k per $k = 1, \dots, K$. Per ciascuna di queste angolazioni, il detector campiona l'intensità del segnale per ognuno dei $P > 0$ raggi X paralleli emessi dalla sorgente, i quali vengono indicati con s_l per $l = 1, \dots, P$. Poiché ogni singola misurazione ottenuta dal detector corrisponde biunivocamente a un raggio che ha attraversato l'oggetto, risulta vantaggioso riorganizzare l'intero set di dati in un'unica sequenza ordinata.

Si definisce quindi un indice globale i , variabile tra 1 e $N = K \cdot P$, che associa ogni coppia di parametri (φ_k, s_l) a un raggio X specifico. In questo modo il percorso dell' i -esimo raggio viene indicato sinteticamente come $r_i(t) := L_{\varphi_k, s_l}(t)$ per ogni $t \in \mathbb{R}$. Grazie a questa indicizzazione il problema inverso può essere inizialmente espresso come un vettore di integrali di linea, in cui ogni componente rappresenta la trasformata di Radon della funzione $f(x_1, x_2)$, che descrive la distribuzione dell'attenuazione dell'oggetto, calcolata lungo la traiettoria r_i :

$$m = \begin{bmatrix} \int_{\mathbb{R}} f(r_1(t)) dt \\ \vdots \\ \int_{\mathbb{R}} f(r_N(t)) dt \end{bmatrix} + \varepsilon$$

Dove $m \in \mathbb{R}^N$ costituisce il vettore delle misurazioni osservate, mentre il termine ε rappresenta il rumore di misurazione additivo presente nel processo.

Il passo successivo riguarda la scomposizione della funzione $f(x_1, x_2)$ in termini discreti. Il dominio bidimensionale di f viene suddiviso in una griglia regolare composta da $M > 0$ pixel, assumendo che il coefficiente di attenuazione mantenga un valore costante all'interno di ciascuno di essi. L'interno dell'oggetto viene così rappresentato come un vettore positivo $\mathbf{f} \in \mathbb{R}^M$, le cui componenti $\mathbf{f}_j$ esprimono il valore del coefficiente del j -esimo pixel.

Tale rappresentazione permette di approssimare la trasformata di Radon di ogni raggio r_i attraverso una sommatoria dei contributi dei singoli elementi

della griglia incontrati lungo il percorso. La singola misurazione m_i viene dunque espressa come

$$m_i = \int_{\mathbb{R}} f(r_i(t)) dt \approx \sum_{j=1}^M a_{i,j} \mathbf{f}_j, \quad \forall i = 1, \dots, N,$$

dove il peso $a_{i,j}$ identifica la lunghezza del tratto percorso dall' i -esimo raggio all'interno del pixel j . Dalla geometria del sistema si deduce che la matrice dell'operatore di forward risultante è caratterizzata da un'elevata sparsità, poiché ogni raggio X interagisce solo con una minima frazione dei pixel che compongono l'immagine.

Il problema della ricostruzione si riduce quindi al sistema lineare

$$A\mathbf{f} + \varepsilon = m, \tag{1.2}$$

dove $A = (a_{i,j}) \in \mathbb{R}^{N \times M}$ è la matrice di proiezione, detta anche operatore di forward. A causa delle grandi dimensioni del sistema e del numero spesso limitato di proiezioni disponibili, il problema risulta frequentemente sotto determinato e mal condizionato, rendendo necessario il ricorso a tecniche di regolarizzazione [9].

Una volta ottenuto il modello discreto (1.2), la ricostruzione tomografica può essere affrontata con diverse classi di metodi. Storicamente, il metodo più diffuso in ambito clinico è il filtered back-projection (FBP), il quale sfrutta la teoria della trasformata di Radon per invertire in modo esplicito l'operatore A mediante operazioni di filtraggio nel dominio della frequenza seguite da una retro proiezione delle misure sull'immagine [6]. Metodi più recenti si basano invece su formulazioni variazionali e algoritmi iterativi di ottimizzazione, i quali introducono esplicitamente un termine di regolarizzazione per ottenere soluzioni stabili anche in presenza di pochi dati o di un elevato livello di rumore [9].

Negli ultimi anni, si sono affermati anche approcci basati sul Machine Learning, che combinano il modello fisico di acquisizione con reti neurali addestrate a migliorare la qualità della ricostruzione. In questo contesto si inseriscono sia i metodi supervisionati, che richiedono grandi dataset di

riferimento, sia metodi non supervisionati come quelli che si basano sull'approccio Deep Image Prior [1]. Questi ultimi hanno il vantaggio di non necessitare di un dataset per l'addestramento, rendendoli particolarmente adatti all'ambito della Tomografia Computerizzata, dove ottenere immagini ground truth affidabili richiederebbe di sottoporre il paziente a dosi di radiazioni inaccettabilmente elevate.

Capitolo 2

Deep Image Prior e Modello di Ottimizzazione Vincolato

In questo capitolo viene presentato il metodo di ricostruzione tomografica basato su reti neurali oggetto di questo lavoro di tesi. Nella Sezione 2.1 viene introdotto l'approccio Deep Image Prior, partendo dalle motivazioni e dall'intuizione alla base del suo sviluppo, per poi descriverne il modello matematico e approfondire il ruolo del prior implicito codificato dall'architettura della rete generatrice. Partendo dalle limitazioni di tale approccio, in particolare dalla dipendenza critica dall'early stopping, si motiva il passaggio a un modello di ottimizzazione vincolato con regolarizzazione automatica. Gli strumenti teorici necessari alla comprensione di questo modello, ovvero la Lagrangiana aumentata e l'operatore prossimale, vengono presentati nella Sezione 2.2 assieme al metodo PGDA utilizzato per la sua risoluzione. Infine, nella Sezione 2.3 viene descritto in dettaglio il modello cDIP-RED proposto da Cascarano et al. [4], che realizza l'approccio di ottimizzazione vincolata con regolarizzazione automatica.

2.1 Deep Image Prior

Per comprendere l'approccio Deep Image Prior (DIP), è utile richiamare il concetto di *prior*. Tale concetto trova le sue radici nella *statistica bayesiana* applicata al *Machine Learning* [5]. La prospettiva bayesiana utilizza la probabilità per esprimere il grado di incertezza di una certa conoscenza. Pertanto, i parametri θ di un modello statistico, che risultano sconosciuti o incerti, vengono rappresentati come una variabile aleatoria. Prima di osservare i dati, rappresentiamo la nostra conoscenza di θ utilizzando la sua distribuzione di probabilità a priori, $p(\theta)$, che chiamiamo *prior* [5]. L'effetto del prior è quello di influenzare il modello durante l'allenamento spostando massa di probabilità verso le regioni dello spazio dei parametri che erano già preferite a priori [5].

Il prior quindi esprime implicitamente le assunzioni iniziali fatte sulla soluzione, e così lo intenderemo da ora in avanti. Inoltre, la sua importanza è quella di orientare il modello verso configurazioni che generano ricostruzioni plausibili. Ad esempio, nel contesto della ricostruzione tomografica, un buon prior favorisce la generazione di immagini che presentano caratteristiche statistiche tipiche delle immagini naturali, scartando al contempo soluzioni rumorose o inconsistenti.

Tradizionalmente, si assume che la capacità di generare immagini realistiche venga appresa interamente dai dati durante la fase di addestramento su grandi dataset. In particolare, si riteneva che le ConvNets eccellessero in questi compiti proprio per la loro abilità nell'estrarre prior complessi delle immagini naturali dai dati stessi [7]. Tuttavia, Zhang et al. [10] hanno dimostrato che reti di classificazione profonde, capaci di generalizzare in modo eccellente su dati reali, sono in grado di adattarsi perfettamente anche a dataset le cui etichette vengono assegnate in modo puramente casuale. Ciò implica che la capacità espressiva di queste reti è tale da consentire la memorizzazione pura dei dati, indipendentemente da qualsiasi struttura statistica.

Partendo da questa osservazione, Ulyanov et al. [7] hanno investigato la

possibilità che la generalizzazione non derivi esclusivamente dall'apprendimento, bensì dalla struttura della rete stessa, mostrando sperimentalmente che l'architettura di una rete generativa convoluzionale è di per sé sufficiente a catturare buona parte delle statistiche delle immagini naturali, ancor prima di qualsiasi addestramento. Questa proprietà strutturale costituisce il prior implicito del metodo Deep Image Prior.

Alla luce di queste premesse, è possibile introdurre formalmente il metodo DIP seguendo la formulazione di Ulyanov et al. [7]. Un generico problema di ricostruzione di immagini ha l'obiettivo di recuperare un'immagine incognita $x \in \mathbb{R}^n$ partendo da una sua misurazione degradata (ad esempio sfocata o rumorosa) $x_0 \in \mathbb{R}^m$. Questo processo può essere modellizzato come un problema lineare inverso:

$$\text{trova } x \in \mathbb{R}^n \quad t.c. \quad Ax + \epsilon = x_0$$

dove il vettore $\epsilon \in \mathbb{R}^m$ denota l'errore additivo proveniente dal rumore di misurazione, mentre l'operatore lineare $A \in \mathbb{R}^{m \times n}$ viene chiamato *operatore di forward* e descrive concettualmente il processo fisico di degradazione o acquisizione [9].

I problemi lineari inversi nell'imaging sono intrinsecamente malposti [4]. Pertanto, per poterne calcolare una soluzione stabile è necessario ricorrere a metodi di regolarizzazione variazionale. Questo comporta la risoluzione di un problema di ottimizzazione nella forma:

$$x^* = \underset{x \in \mathbb{R}^n}{\operatorname{argmin}} \left\{ \frac{1}{2} \|Ax - x_0\|_2^2 + \lambda R(x) \right\} \quad (2.1)$$

dove il primo termine rappresenta la *data fidelity* (o fedeltà dei dati), mentre il secondo viene chiamato termine di regolarizzazione. Quest'ultimo è composto dal parametro di regolarizzazione $\lambda > 0$ e dal funzionale di regolarizzazione $R: \mathbb{R}^n \rightarrow \mathbb{R}$. Il funzionale $R(x)$ ha il compito di catturare un generico prior riguardo le immagini naturali, mentre il parametro di regolarizzazione pesa l'importanza da assegnare alla regolarizzazione rispetto all'aderenza ai dati misurati [4]. La definizione del termine regolarizzato-

re è storicamente complessa e dipendente dal dominio. Inoltre, la selezione ottimale del parametro λ richiede spesso svariate prove empiriche [4].

A questo punto, nel framework DIP, si introduce l'utilizzo dell'architettura di una rete ConvNet generatrice per sostituire la regolarizzazione esplicita. Sia $f_\theta: \mathbb{R}^N \rightarrow \mathbb{R}^n$ una ConvNet generatrice parametrizzata dai propri pesi $\theta \in \mathbb{R}^s$, che, dato un vettore input $z \in \mathbb{R}^N$, calcola l'immagine $x = f_\theta(z)$. La rete neurale f_θ agisce quindi come una parametrizzazione dell'immagine stessa. Piuttosto che cercare l'immagine x nello spazio non vincolato $\mathbb{R}^n$ applicando un *prior* matematico mediante l'uso di regolarizzazione esplicita, l'approccio DIP omette completamente il termine $\lambda R(x)$ in (2.1) e lo sostituisce forzando l'immagine a essere generata dalla rete. Il problema di ottimizzazione viene riformulato nello spazio dei parametri della rete:

$$\theta^* = \operatorname{argmin}_{\theta \in \mathbb{R}^s} \frac{1}{2} \|A f_\theta(z) - x_0\|_2^2, \quad x^* = f_{\theta^*}(z) \quad (2.2)$$

dove il minimizzatore θ^* viene calcolato utilizzando un metodo di ottimizzazione (ad esempio il metodo di discesa del gradiente) partendo da un'inizializzazione casuale dei pesi θ . In questa formulazione, l'immagine prior codificato dall'architettura di f_θ impone naturalmente una forte regolarità alla soluzione. Questo avviene poiché lo spazio delle possibili soluzioni viene implicitamente ristretto all'insieme delle sole immagini che la rete è strutturalmente in grado di generare, ovvero immagini con caratteristiche statistiche tipiche delle immagini naturali [7].

Un'ultima ma fondamentale considerazione riguarda la risoluzione pratica del problema. Come discusso in precedenza, gli studi di Zhang et al. [10] hanno dimostrato che l'elevata capacità espressiva della ConvNets consente loro di memorizzare qualsiasi tipo di dato, comprese immagini composte da puro rumore casuale. Pertanto esistono sicuramente dei pesi θ per cui $f_\theta(z)$ ricostruisce immagini completamente rumorose, e quindi è lecito domandarsi in che modo la parametrizzazione discussa in precedenza offre un restringimento dello spazio delle soluzioni.

Sebbene questo sia vero, Ulyanov et al. [7] mostrano come l'architettura della rete influenzi positivamente la ricerca della soluzione durante l'otti-

mizzazione. In particolare, osservano che la rete oppone una forte inerzia alle soluzioni "cattive" mentre favorisce la discesa verso immagini naturali. In altre parole la parametrizzazione offre alta impedenza al rumore e bassa impedenza al segnale [7].

Per sfruttare questa proprietà, nella sua formulazione originale il metodo DIP viene accoppiato a tecniche di *early stopping*. Interrompendo prematuramente l'ottimizzazione, si restringe di fatto lo spazio delle soluzioni alle sole immagini generabili da $f_\theta(z)$ con dei pesi θ che non discostano troppo dalla loro iniziale distribuzione casuale [7].

Tuttavia, la dipendenza dall'early stopping introduce una sensibilità critica al numero di iterazioni, la cui scelta ottimale risulta difficile da determinare a priori. Questo limite ha motivato lo sviluppo di modelli DIP con regolarizzazione automatica, come il modello cDIP-RED proposto da Cascarano et al. [4], che viene discusso di seguito applicandolo alla ricostruzione tomografica.

Partendo dal problema (1.2) mostrato nella Sezione 1.1, viene adottata la seguente notazione:

$$Ax + \varepsilon = x_0,$$

dove $x_0 \in \mathbb{R}^N$ è il sinogramma osservato su un numero di $N = K \cdot P$, con $K > 0$ numero di angoli e P numero di rilevazioni per angolo, $x \in \mathbb{R}^M$ è l'immagine vettorizzata dell'oggetto osservato, i cui pixel esprimono il coefficiente di attenuazione in quella posizione, $A \in \mathbb{R}^{N \times M}$ è il forward operator definito nella Sezione 1.1, mentre $\varepsilon \in \mathbb{R}^N$ è il rumore della misurazione, che in questo lavoro viene modellizzato come Rumore Bianco Additivo Gaussiano (AWGN) con media zero e deviazione standard σ_ε .

Siccome il lavoro è condotto in un contesto simulato, l'immagine da ricostruire x_{real} è nota con una risoluzione virtualmente infinita, e viene utilizzata per simulare la misurazione rumorosa x_0 . Bisogna quindi prestare attenzione nel calcolare i dati della ricostruzione x_0 e A in maniera tale da non costituire un *inverse crime*, ovvero la situazione in cui l'ottima qualità della ricostruzione è dovuta all'utilizzo dello stesso modello sia per la simulazione dei dati che per la ricostruzione [9]. Quindi il sinogramma rumoroso x_0 non

viene calcolato applicando direttamente l'operatore A all'oggetto x_{real} , bensì viene calcolato applicando un operatore di forward $A' \in \mathbb{R}^{N \times M'}$ a una versione dell'oggetto x_{real} campionata a una risoluzione doppia $x_{\text{fine}} \in \mathbb{R}^{M'}$, con $M' = 4M$:

$$x_0 = A' x_{\text{fine}} + \varepsilon. \quad (2.3)$$

In questo modo, il dato acquisito è prodotto da un operatore discreto diverso da quello usato nella ricostruzione, garantendo una separazione netta tra i due modelli [9]. L'operatore A utilizzato nella fase di ricostruzione viene quindi costruito sulla griglia di risoluzione target M .

Sia inoltre $f_\theta: \mathbb{R}^M \rightarrow \mathbb{R}^M$ una ConvNet generatrice parametrizzata dai pesi $\theta \in \mathbb{R}^s$, che dato un vettore di input $z \in \mathbb{R}^M$ produce l'immagine $x = f_\theta(z)$. L'approccio DIP classico presentato nella Sezione 2.1 consiste nel trovare θ^* risolvendo il problema (2.2) tramite discesa del gradiente combinata con early stopping [7], ottenendo l'immagine ricostruita $x^* = f_{\theta^*}(z)$.

Per superare la dipendenza critica dal numero di iterazioni introdotta dall'early stopping, negli ultimi anni sono stati sviluppati metodi che affiancano al termine di fedeltà ai dati un esplicito termine di regolarizzazione [4]. Il problema di ottimizzazione diventa quindi:

$$\theta^* = \underset{\theta \in \mathbb{R}^s}{\operatorname{argmin}} \left\{ \frac{1}{2} \|A f_\theta(z) - x_0\|_2^2 + \lambda R(f_\theta(z)) \right\}, \quad (2.4)$$

dove $R: \mathbb{R}^M \rightarrow \mathbb{R}$ è una funzione di regolarizzazione, mentre $\lambda > 0$ è il parametro di regolarizzazione che pesa il contributo di R rispetto alla fedeltà ai dati. La scelta ottimale di λ rimane però un problema aperto, che nel caso dei metodi DIP si risolve spesso con la regolazione manuale del parametro tramite una procedura *trial-and-error* [4].

Per ovviare a questo problema, Cascarano et al. [4] propongono due metodi DIP per l'ottimizzazione di (2.4) che eliminano la necessità di calcolare manualmente il parametro di regolarizzazione λ .

Il primo, denominato DIP-WTV, è un modello non vincolato che adotta una versione spazialmente adattiva della Variazione Totale (TV), stimando

automaticamente i parametri di regolarizzazione locali durante l'ottimizzazione. La Variazione Totale si basa sull'osservazione che le immagini naturali sono tipicamente costanti a tratti, in quanto presentano spesso regioni uniformi separate da bordi netti, il che implica che il loro gradiente sia non nullo solo in corrispondenza di tali bordi. Formalmente, la TV di un'immagine vettorizzata $u \in \mathbb{R}^n$ è definita come:

$$\text{TV}(u) = \sum_{i=1}^n \sqrt{(\Delta_h u_i)^2 + (\Delta_v u_i)^2},$$

dove $\Delta_h u_i$ e $\Delta_v u_i$ sono le differenze finite del primo ordine del pixel i rispettivamente lungo la direzione orizzontale e verticale [4].

Il secondo, denominato cDIP-RED, è un modello vincolato che utilizza un regolarizzatore RED (*Regularization by Denoising*). L'idea alla base è che un denoiser $D(\cdot)$ possa essere utilizzato per indurre regolarizzazione. Infatti, se un'immagine u presenta poco rumore, allora il denoiser non la modifica sostanzialmente e quindi $D(u) \approx u$, mentre se u presenta molto rumore, allora la differenza $u - D(u)$ risulta particolarmente grande. La funzione di regolarizzazione è quindi definita come [8]:

$$R(u) = \frac{1}{2} u^T (u - D(u)). \quad (2.5)$$

Aggiungere $R(u)$ come regolarizzazione spinge quindi l'ottimizzazione verso immagini con poco rumore, penalizzando implicitamente quelle che ne presentano in quantità elevata. Un ulteriore vantaggio riguarda il calcolo del suo gradiente rispetto a u , il quale è semplicemente dato da:

$$\nabla_u R(u) = u - D(u), \quad (2.6)$$

senza dover differenziare il denoiser $D(\cdot)$ [8].

Nel metodo cDIP-RED, quindi, la forza della regolarizzazione viene determinata automaticamente imponendo, secondo il principio di discrepanza di Morozov, che il residuo $\|Af_\theta(z) - x_0\|_2^2$ rimanga coerente con il livello di rumore presente nei dati. In questo modo si elimina completamente la necessità di fissare un valore di λ a priori.

Entrambi i metodi vengono formulati come problemi minimax e risolti attraverso una versione modificata del metodo *Proximal Gradient Descent-Ascent* (PGDA) proposta da Cascarano et al. [4]. Gli strumenti teorici necessari per la comprensione di tale approccio, in particolare la Lagrangiana aumentata e l'operatore di prossimità, vengono introdotti nella Sezione 2.2 assieme al metodo PGDA stesso. Il modello cDIP-RED, oggetto di questo lavoro di tesi, viene descritto in dettaglio nella Sezione 2.3.

2.2 Metodo PGDA

Per introdurre in maniera naturale il metodo PGDA, è opportuno partire dalla nozioni di Lagrangiana, mostrando come essa trasformi un problema vincolato in un problema di tipo minimax, al quale il metodo è direttamente applicabile.

Il problema di ottimizzazione vincolato al quale si fa riferimento in questa sezione presenta la seguente forma:

$$\begin{aligned} \min_{x \in \mathbb{R}^n} f(x) \quad \text{t.c.} \quad & f_i(x) \leq 0, \quad i = 1, \dots, m \\ & h_i(x) = 0, \quad i = 1, \dots, p, \end{aligned} \tag{2.7}$$

dove $f, f_1, \dots, f_m: \mathbb{R}^n \rightarrow \mathbb{R}$ sono le funzioni obiettivo e di vincolo di disuguaglianza, mentre $h_1, \dots, h_p: \mathbb{R}^n \rightarrow \mathbb{R}$ descrivono i vincoli di uguaglianza del problema. Si denota con p^* il valore ottimo del problema.

Definizione 2.2.1 (Lagrangiana). *La funzione Lagrangiana associata al problema vincolato (2.7) è la funzione $\mathcal{L}: \mathbb{R}^n \times \mathbb{R}^m \times \mathbb{R}^p \rightarrow \mathbb{R}$ definita come:*

$$\mathcal{L}(x, \lambda, \nu) = f(x) + \sum_{i=1}^m \lambda_i f_i(x) + \sum_{i=1}^p \nu_i h_i(x)$$

dove $\lambda \in \mathbb{R}^m$ e $\nu \in \mathbb{R}^p$ sono detti moltiplicatori di Lagrange associati rispettivamente ai vincoli di disuguaglianza e di uguaglianza.

L'idea fondamentale della Lagrangiana consiste nell'incorporare i vincoli all'interno della funzione obiettivo tramite i moltiplicatori di Lagrange, sostituendo i vincoli espliciti con una penalizzazione lineare. Il valore di ciascun

moltiplicatore misura quanto sia influente il vincolo corrispondente: un valore elevato indica che una piccola perturbazione del vincolo avrebbe un impatto significativo sulla soluzione ottima, mentre un valore ridotto segnala che il vincolo può essere rilassato senza alterarne significativamente il risultato [2].

Ponendo la condizione $\lambda \geq 0$, il problema originale (2.7) può essere ricondotto al seguente problema di minimax [2]:

$$\min_{x \in \mathbb{R}^n} \max_{\substack{\lambda \in \mathbb{R}_+^m \\ \nu \in \mathbb{R}^p}} \mathcal{L}(x, \lambda, \nu).$$

Fissato un punto x , massimizzare $L(x, \lambda, \nu)$ rispetto ai moltiplicatori fa divergere il valore a $+\infty$ se x viola anche un solo vincolo, mentre restituisce $f(x)$ se x rispetta tutti i vincoli. La minimizzazione esterna scarta quindi tutti i punti x non ammissibili, in quanto associati a un valore infinito, e tra quelli ammissibili minimizza $f(x)$, coincidendo così con il problema originale.

Per poter descrivere formalmente lo schema iterativo del metodo PGDA è necessario introdurre un'ultima nozione, quella di operatore prossimale.

Definizione 2.2.2 (Operatore prossimale). *Dato un $\alpha > 0$ qualsiasi e una funzione $f: \mathbb{R}^d \rightarrow \mathbb{R} \cup \{+\infty\}$ semi-continua inferiormente e convessa, l'operatore prossimale di f di parametro α è la funzione $\text{prox}_{\alpha f}: \mathbb{R}^d \rightarrow \mathbb{R}^d$ definita come [3]:*

$$\text{prox}_{\alpha f}(x) := \underset{z \in \mathbb{R}^d}{\operatorname{argmin}} \left\{ f(z) + \frac{1}{2\alpha} \|z - x\|^2 \right\}$$

Intuitivamente, $\text{prox}_{\alpha f}(x)$ restituisce il punto che bilancia la minimizzazione di f con la vicinanza al punto corrente x . Il parametro α controlla il peso relativo dei due obiettivi.

Il metodo PGDA è stato proposto per risolvere problemi di minimax della forma [3]:

$$\min_{x \in \mathbb{R}^d} \max_{y \in \mathbb{R}^n} \{ \Psi(x, y) := \Phi(x, y) + \mathcal{R}(x) + h(y) \} \quad (2.8)$$

dove la funzione di accoppiamento $\Phi: \mathbb{R}^d \times \mathbb{R}^n \rightarrow \mathbb{R}$ è debolmente convessa in x e concava in y , mentre $\mathcal{R}$ e h sono funzioni proprie, convesse e semi-continue inferiormente che fungono da regolarizzatori rispettivamente su x e su y .

Partendo da un punto iniziale $(x^0, y^0) \in \mathbb{R}^d \times \mathbb{R}^n$, il metodo PGDA risolve il problema (2.8) applicando il seguente schema iterativo:

$$\begin{cases} x^{k+1} = \text{prox}_{\alpha_x \mathcal{R}}(x^k - \alpha_x \nabla_x \Phi(x^k, y^k)) \\ y^{k+1} = \text{prox}_{\alpha_y h}(y^k + \alpha_y \nabla_y \Phi(x^{k+1}, y^k)) \end{cases}, \quad (2.9)$$

dove $\alpha_x > 0$ e $\alpha_y > 0$ sono learning rates dei passi di aggiornamento per le due variabili. Il passo su x esegue una discesa sul gradiente prossimale, mentre il passo su y esegue un'ascesa, in accordo con la struttura minimax del problema.

Nel contesto del metodo cDIP-RED, che verrà illustrato nella sezione successiva, il regolarizzatore h è assente, ovvero $h \equiv 0$. Applicando la Definizione 2.2.2, il passo su y si semplifica in un puro passo di ascesa sul gradiente, e lo schema (2.9) diventa:

$$\begin{cases} x^{k+1} = \text{prox}_{\alpha_x \mathcal{R}}(x^k - \alpha_x \nabla_x \Phi(x^k, y^k)) \\ y^{k+1} = y^k + \alpha_y \nabla_y \Phi(x^{k+1}, y^k) \end{cases}. \quad (2.10)$$

2.3 Modello cDIP-RED

Il modello cDIP-RED di Cascarano et al. [4] affronta il problema regolarizzato (2.4) riformulandolo come un problema di ottimizzazione vincolata, sfruttando il principio di discrepanza di Morozov. Tale principio afferma che una ricostruzione è da considerarsi accettabile quando il residuo sui dati è coerente con il livello di rumore presente nelle misurazioni, ovvero quando vale [9]:

$$\|Af_\theta(z) - x_0\| \leq \|\varepsilon\|,$$

siccome $\varepsilon \sim \mathcal{N}(0, \sigma_\varepsilon^2)$, come viene assunto nella Sezione 2.1, allora si ha che $\mathbb{E}[\|\varepsilon\|] = \sqrt{N}\sigma_\varepsilon$ [9], quindi il vincolo di discrepanza viene posto rispetto a $\delta = \sqrt{N}\sigma_\varepsilon$.

Quindi, adottando il principio di discrepanza, il problema (2.4) viene riformulato come [4]:

$$\operatorname{argmin}_{\theta \in \mathbb{R}^s} R(f_\theta(z)) \quad \text{t.c.} \quad f_\theta(z) \in \mathcal{D}_{\sigma_\varepsilon}, \quad (2.11)$$

dove R è una funzione di regolarizzazione RED definita come in (2.5), mentre $\mathcal{D}_{\sigma_\varepsilon} \subseteq \mathbb{R}^M$ è definito come:

$$\mathcal{D}_{\sigma_\varepsilon} := \{f_\theta(z) \mid \|Af_\theta(z) - x_0\|_2^2 \leq \tau\sigma_\varepsilon^2 N\},$$

con $\tau > 0$ parametro di tolleranza che assorbe la variabilità statistica di $\delta^2 = N\sigma_\varepsilon^2$.

Cascarano et al. [4] dimostrano che, sotto opportune ipotesi di regolarità, il problema vincolato (2.11) e il problema regolarizzato (2.4) ammettono la stessa soluzione, rendendo la formulazione vincolata equivalente a quella regolarizzata.

Il vantaggio principale di questa riformulazione è l'eliminazione del parametro di regolarizzazione λ : il vincolo di discrepanza assume il suo ruolo, determinando automaticamente la forza della regolarizzazione in base al livello di rumore nei dati. L'unica quantità da stimare diventa quindi σ_ε che, pur non essendo nota a priori, può essere approssimata con gli strumenti moderni di stima del rumore in modo molto più agevole rispetto alla calibrazione manuale di λ [4].

Per la risoluzione di (2.11), vengono introdotte due variabili ausiliarie:

$$\begin{aligned} t &:= f_\theta(z), \\ r &:= Af_\theta(z) - x_0. \end{aligned}$$

Introducendo tali variabili, il problema (2.11) viene riscritto come:

$$\operatorname{argmin}_{\substack{\theta \in \mathbb{R}^s \\ t \in \mathbb{R}^M, r \in \mathbb{R}^N}} \{R(t) + i_{B_\rho}(r)\} \quad \text{t.c.} \quad t = f_\theta(z) \text{ e } r = Af_\theta(z) - x_0,$$

dove $B_\rho := \{r \in \mathbb{R}^N \mid \|r\|_2 \leq \rho\}$ è la sfera centrata nell'origine di raggio $\rho := \sqrt{\tau\sigma_\varepsilon^2 N}$, e $i_{B_\rho}: \mathbb{R}^N \rightarrow \mathbb{R} \cup \{+\infty\}$ è la funzione indicatrice di B_ρ , definita

come:

$$i_{B_\rho}(r) := \begin{cases} 0 & \text{se } \|r\|_2 \leq \rho \\ +\infty & \text{altrimenti.} \end{cases}$$

Si osserva che la presenza di $i_{B_\rho}(r)$ nella funzione obiettivo impone implicitamente il vincolo di discrepanza, infatti $i_{B_\rho}(r) = +\infty$ per ogni $r \notin B_\rho$ e quindi qualsiasi soluzione con $\|Af_\theta(z) - x_0\|_2 > \rho$ viene automaticamente esclusa dalla minimizzazione.

Applicando la Definizione 2.2.1 si ottiene la Lagrangiana:

$$\begin{aligned} \mathcal{L}_0(\theta, t, r; \mu_t, \mu_r) = & R(t) + i_{B_\rho}(r) + \langle \mu_t, f_\theta(z) - t \rangle \\ & + \langle \mu_r, Af_\theta(z) - x_0 - r \rangle, \end{aligned}$$

con $\mu_t \in \mathbb{R}^M$ e $\mu_r \in \mathbb{R}^N$ moltiplicatori di Lagrange relativi alle rispettive variabili. Cascarano et al. [4] adottano però la forma aumentata della Lagrangiana, che introduce due parametri di penalizzazione positivi β_t e β_r per migliorare la convergenza dell'ottimizzazione:

$$\begin{aligned} \mathcal{L}(\theta, t, r; \mu_t, \mu_r) = & \mathcal{L}_0 + \frac{\beta_t}{2} \|t - f_\theta(z)\|_2^2 \\ & + \frac{\beta_r}{2} \|r - (Af_\theta(z) - x_0)\|_2^2 \end{aligned} \quad (2.12)$$

Per ricondurre (2.12) alla forma minimax (2.8) richiesta dal metodo PGDA, si introducono le variabili $x := [\theta; t; r] \in \mathbb{R}^{s+M+N}$ e $y := [\mu_t; \mu_r] \in \mathbb{R}^{M+N}$, e si decompone la Lagrangiana aumentata nelle due componenti:

$$\begin{aligned} K(x, y) = & \frac{\beta_t}{2} \|t - f_\theta(z)\|_2^2 + \frac{\beta_r}{2} \|r - (Af_\theta(z) - x_0)\|_2^2 \\ & + \langle \mu_t, f_\theta(z) - t \rangle + \langle \mu_r, Af_\theta(z) - x_0 - r \rangle, \end{aligned} \quad (2.13)$$

e $\mathcal{R}(x) = R(t) + i_{B_\rho}(r)$, in modo che $\mathcal{L}(x, y) = K(x, y) + \mathcal{R}(x)$. Il problema da risolvere diventa quindi:

$$\min_{x \in \mathbb{R}^{s+M+N}} \max_{y \in \mathbb{R}^{M+N}} \mathcal{L}(x, y),$$

che corrisponde esattamente a (2.8) con $h \equiv 0$. Si applica quindi lo schema iterativo (2.10), che in questo caso diventa:

$$\begin{cases} x^{k+1} = \text{prox}_{\alpha_x \mathcal{R}}(x^k - \alpha_x \nabla_x K(x^k, y^k)) \\ y^{k+1} = y^k + \alpha_y \nabla_y K(x^{k+1}, y^k), \end{cases} \quad (2.14)$$

con α_x e α_y learning rates positivi.

Il modello cDIP-RED viene implementato calcolando l'iterazione (2.14), che tuttavia non è ancora espressa in una forma direttamente implementabile. Di seguito si ricava quindi un'espressione esplicita per ciascuno dei passi dell'iterazione.

Il passo sui moltiplicatori y è di immediata risoluzione. Siccome $y = [\mu_t; \mu_r]$, il gradiente $\nabla_y K(x^{k+1}, y^k)$ si separa nelle due componenti:

$$\begin{aligned}\nabla_{\mu_t} K(x^{k+1}, y^k) &= f_{\theta^{k+1}}(z) - t^{k+1}, \\ \nabla_{\mu_r} K(x^{k+1}, y^k) &= Af_{\theta^{k+1}}(z) - x_0 - r^{k+1}.\end{aligned}$$

Il passo di aggiornamento di y in (2.14) diventa quindi:

$$\begin{cases} \mu_t^{k+1} = \mu_t^k + \alpha_y (f_{\theta^{k+1}}(z) - t^{k+1}) \\ \mu_r^{k+1} = \mu_r^k + \alpha_y (Af_{\theta^{k+1}}(z) - x_0 - r^{k+1}) \end{cases} \quad (2.16)$$

Il passo sulle variabili x risulta invece più complesso. Espandendo l'operatore prossimale secondo la Definizione 2.2.2, il passo su x in (2.14) diventa:

$$x^{k+1} = \underset{x \in \mathbb{R}^{s+M+N}}{\operatorname{argmin}} \left\{ \alpha_x \mathcal{R}(x) + \frac{1}{2} \|x - (x^k - \alpha_x \nabla_x K(x^k, y^k))\|_2^2 \right\}. \quad (2.17)$$

Separando le componenti di $x = [\theta; t; r]$ e sfruttando l'additività della norma al quadrato per vettori a blocchi, il passo (2.17) diventa:

$$\begin{aligned} x^{k+1} &= \underset{[\theta; t; r]}{\operatorname{argmin}} \left\{ \alpha_x \mathcal{R}([\theta; t; r]) + \frac{1}{2} \left\| \begin{bmatrix} \theta \\ t \\ r \end{bmatrix} - \begin{bmatrix} \theta^k - \alpha_x \nabla_{\theta} K(x^k, y^k) \\ t^k - \alpha_x \nabla_t K(x^k, y^k) \\ r^k - \alpha_x \nabla_r K(x^k, y^k) \end{bmatrix} \right\|_2^2 \right\} \\ &= \underset{[\theta; t; r]}{\operatorname{argmin}} \{ \alpha_x R(t) + \alpha_x i_{B_\rho}(r) \\ &\quad + \frac{1}{2} \|\theta - (\theta^k - \alpha_x \nabla_{\theta} K(x^k, y^k))\|_2^2 \\ &\quad + \frac{1}{2} \|t - (t^k - \alpha_x \nabla_t K(x^k, y^k))\|_2^2 \\ &\quad + \frac{1}{2} \|r - (r^k - \alpha_x \nabla_r K(x^k, y^k))\|_2^2 \}. \end{aligned} \quad (2.18)$$

Siccome l'obiettivo in (2.18) è separabile rispetto alle variabili θ , t e r , il passo si riduce a tre sotto-problemi indipendenti, uno per ciascuna componente:

$$\begin{cases} \theta^{k+1} = \operatorname{argmin}_{\theta \in \mathbb{R}^s} \frac{1}{2} \|\theta - (\theta^k - \alpha_x \nabla_{\theta} K(x^k, y^k))\|_2^2 \\ t^{k+1} = \operatorname{argmin}_{t \in \mathbb{R}^M} \left\{ \alpha_x R(t) + \frac{1}{2} \|t - (t^k - \alpha_x \nabla_t K(x^k, y^k))\|_2^2 \right\} \\ r^{k+1} = \operatorname{argmin}_{r \in \mathbb{R}^N} \left\{ \alpha_x i_{B_\rho}(r) + \frac{1}{2} \|r - (r^k - \alpha_x \nabla_r K(x^k, y^k))\|_2^2 \right\} \end{cases} \quad (2.19)$$

I sotto-problemi in (2.19) non ammettono ancora una soluzione esplicita immediata. Di seguito viene quindi ricavata la forma chiusa per ciascuno di essi.

Aggiornamento di θ

Il primo sotto-problema riguarda l'aggiornamento dei pesi della rete f_{θ} , che si risolve direttamente con un passo di discesa del gradiente:

$$\theta^{k+1} = \theta^k - \alpha_x \nabla_{\theta} K(x^k, y^k).$$

È quindi necessario calcolare il gradiente di K rispetto ai pesi della rete. Da (2.13) si osserva che K dipende da θ esclusivamente attraverso l'output della rete $f_{\theta}(z)$, quindi per la regola della catena [5]:

$$\nabla_{\theta} K = \left(\frac{\partial f_{\theta}}{\partial \theta} \right) \nabla_f K,$$

dove $\frac{\partial f_{\theta}}{\partial \theta}$ denota la derivata parziale della rete rispetto a tutti i suoi parametri. Per reti CNN con milioni di parametri, questo calcolo non è affrontabile direttamente in quanto avrebbe costo computazionale e costo di memoria proibitivi. Si ricorre quindi all'algoritmo di *back-propagation* del gradiente [5], implementato nelle principali librerie di deep learning. Rimane pertanto da calcolare $\nabla_f K$, che differenziando K rispetto a $f = f_{\theta}(z)$ risulta:

$$\nabla_f K(x, y) = \beta_t (f_{\theta}(z) - t) + \mu_t + A^{\top} (\beta_r (A f_{\theta}(z) - x_0 - r) + \mu_r). \quad (2.20)$$

Concludendo, nella pratica per l'aggiornamento dei pesi θ è sufficiente calcolare (2.20) e passarlo all'algoritmo di back-propagation implementato dalla

libreria di deep learning utilizzata, il quale si occupa di propagarlo all'indietro lungo tutta la rete.

Aggiornamento di t^{k+1}

Successivamente, il secondo sotto-problema riguarda l'aggiornamento della variabile ausiliaria t . Differenziando K rispetto a t si ottiene:

$$\nabla_t K(x^k, y^k) = \beta_t(t^k - f_{\theta^{k+1}}(z)) - \mu_t^k.$$

dove si utilizza $f_{\theta^{k+1}}(z)$ in quanto l'aggiornamento dei parametri θ è già stato eseguito nel passo precedente. Sostituendo nel secondo sotto-problema di (2.19) e raccogliendo rispetto a t^k si ottiene:

$$t^{k+1} = \operatorname{argmin}_{t \in \mathbb{R}^M} \left\{ \alpha_x R(t) + \frac{1}{2} \|t - [(1 - \alpha_x \beta_t)t^k + \alpha_x (\beta_t f_{\theta^{k+1}}(z) + \mu_t)]\|_2^2 \right\}.$$

Definendo la variabile ausiliaria:

$$v_t^k := (1 - \alpha_x \beta_t) t^k + \alpha_x (\beta_t f_{\theta^{k+1}}(z) + \mu_t^k),$$

il sotto-problema si riduce a:

$$t^{k+1} = \operatorname{argmin}_{t \in \mathbb{R}^M} \left\{ \alpha_x R(t) + \frac{1}{2} \|t - v_t^k\|_2^2 \right\}. \quad (2.21)$$

Per risolvere (2.21), si segue la strategia di punto fisso proposta da Mataev et al. in [8] per il modello DeepRED. Sfruttando (2.5), il gradiente della funzione obiettivo è:

$$\nabla_t \left(\alpha_x R(t) + \frac{1}{2} \|t - v_t^k\|_2^2 \right) = \alpha_x (t - \mathbf{D}(t)) + (t - v_t^k).$$

Azzerandolo e raccogliendo i termini in t si ricava l'equazione di punto fisso:

$$t = \frac{\alpha_x \mathbf{D}(t) + v_t^k}{1 + \alpha_x},$$

che si risolve iterativamente partendo da $t^{(0)} = t^k$.

Seguendo [8, 4], nella pratica si esegue una sola iterazione, ottenendo:

$$t^{k+1} = \frac{\alpha_x \mathbf{D}(t^k) + v_t^k}{1 + \alpha_x}. \quad (2.22)$$

Aggiornamento di r^{k+1}

Infine, il terzo sotto-problema riguarda l'aggiornamento della variabile ausiliaria r . Differenziando K rispetto a r si ottiene:

$$\nabla_r K(x^k, y^k) = \beta_r (r^k - (Af_{\theta^{k+1}}(z) - x_0)) - \mu_r^k,$$

dove, come nel caso precedente, si utilizza $f_{\theta^{k+1}}(z)$ in quanto l'aggiornamento dei parametri θ è già stato eseguito. Sostituendo nel terzo sotto-problema di (2.19) e raccogliendo rispetto a r^k , si ottiene:

$$r^k - \alpha_x \nabla_r K(x^k, y^k) = (1 - \alpha_x \beta_r) r^k + \alpha_x (\beta_r (Af_{\theta^{k+1}}(z) - x_0) + \mu_r)$$

Definendo la variabile ausiliaria:

$$v_r^k := (1 - \alpha_x \beta_r) r^k + \alpha_x (\beta_r (Af_{\theta^{k+1}}(z) - x_0) + \mu_r^k),$$

il sotto-problema si riduce a:

$$r^{k+1} = \operatorname{argmin}_{r \in \mathbb{R}^N} \left\{ \alpha_x i_{B_\rho}(r) + \frac{1}{2} \|r - v_r^k\|_2^2 \right\}. \quad (2.23)$$

Siccome $i_{B_\rho}(r) = +\infty$ per ogni $r \notin B_\rho$, la presenza della funzione indicatrice nella funzione obiettivo di (2.23) impone implicitamente il vincolo $r \in B_\rho$, riducendo il problema alla proiezione di v_r^k sulla sfera B_ρ :

$$r^{k+1} = \operatorname{argmin}_{r \in B_\rho} \frac{1}{2} \|r - v_r^k\|_2^2 \quad (2.24)$$

Si osserva che se $v_r^k \in B_\rho$, il minimo è raggiunto in $r^{k+1} = v_r^k$. Se invece $v_r^k \notin B_\rho$, il punto ammissibile più vicino a v_r^k è quello sul bordo della sfera lungo la direzione di v_r^k , ottenuto moltiplicando il suo vettore unitario per il raggio della sfera ρ . La soluzione di (2.24) è quindi:

$$r^{k+1} = \begin{cases} v_r^k & \text{se } \|v_r^k\|_2 \leq \rho, \\ \frac{v_r^k}{\|v_r^k\|_2} \rho & \text{altrimenti.} \end{cases} \quad (2.25)$$

I risultati ottenuti permettono ora di scrivere in forma esplicita un'iterazione completa del metodo cDIP-RED. Ogni iterazione esegue in sequenza l'aggiornamento dei parametri della rete, delle variabili t e r , e infine dei moltiplicatori di Lagrange μ_t e μ_r . L'Algoritmo 1 riassume un passo iterativo del metodo.

Algoritmo 1 Iterazione del metodo cDIP-RED

Input: $\theta^k, t^k, r^k, \mu_t^k, \mu_r^k$ **Output:** $\theta^{k+1}, t^{k+1}, r^{k+1}, \mu_t^{k+1}, \mu_r^{k+1}$

- 1: Calcola $\nabla_f K(x^k, y^k)$ tramite (2.20)
 - 2: Aggiorna θ^{k+1} via back-propagation
 - 3: Calcola v_t^k e aggiorna t^{k+1} tramite (2.22)
 - 4: Calcola v_r^k e aggiorna r^{k+1} tramite (2.25)
 - 5: Aggiorna μ_t^{k+1} e μ_r^{k+1} tramite (2.16)
-

2.4 Dettagli implementativi

In questa sezione vengono concretizzate le scelte che la formulazione del metodo cDIP-RED, presentata nella Sezione 2.3, lascia aperte. Alcune di queste riguardano parametri numerici dell'algoritmo, altre dipendono invece dalla struttura specifica del problema tomografico affrontato in questo lavoro.

Architettura della rete. La rete generatrice f_θ adotta un'architettura Encoder-Decoder ispirata alla U-Net utilizzata da Ulyanov et al. [7]. In questo tipo di architettura, l'encoder riduce progressivamente la risoluzione spaziale dell'input attraverso una serie di livelli convoluzionali, fino a raggiungere una rappresentazione compatta chiamata collo di bottiglia, o bottleneck. Il decoder ricostruisce poi l'immagine ripristinando la risoluzione originale tramite operazioni di upsampling. Inoltre, il numero di canali rimane costante lungo tutti i livelli della rete.

Per individuare la configurazione più adatta al problema, sono state testate diverse combinazioni di profondità e numero di canali, i cui risultati sono riportati in Figura 2.1.

Con immagini di $M = 128$ pixel, una profondità pari a 5 porta il bottleneck a soli 4×4 pixel, una compressione spaziale ritenuta eccessiva. Una profondità pari a 4 produce invece un bottleneck di 8×8 , più adeguato. Tra le configurazioni a profondità 4, quella con 128 canali ottiene la ricostruzione

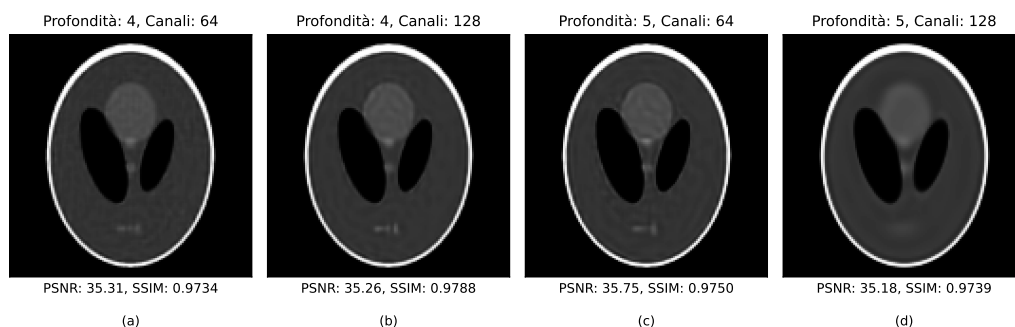


Figura 2.1: Ricostruzioni ottenute dal modello cDIP-RED al variare della profondità della rete (4 e 5 livelli) e del numero di canali per livello (64 e 128). Per ciascuna configurazione sono riportati il PSNR e lo SSIM rispetto al phantom di riferimento. L'immagine (a) corrisponde alla configurazione selezionata.

ritenuta migliore, sia per lo SSIM più elevato tra tutte le varianti (0.9788) sia per la qualità visiva, vedi Figura 2.1 (c).

È stata valutata anche l'introduzione di *skip connections* tra encoder e decoder, nella forma proposta da Baguer et al. [1]. Come mostrato in Figura 2.2, queste non apportano alcun miglioramento apprezzabile, introducendo anzi un leggero peggioramento delle metriche probabilmente dovuto all'aggiunta di rumore nel percorso di ricostruzione.

La configurazione adottata nell'esperimento è quindi quella con profondità 4, 128 canali per livello e nessuna skip connection, per un totale di circa 1.6 milioni di parametri addestrabili.

Regolarizzazione automatica e parametri di penalizzazione. Come discusso nella Sezione 2.3, uno dei vantaggi principali del modello cDIP-RED è l'eliminazione del parametro di regolarizzazione λ , la cui scelta ottimale richiede spesso una calibrazione manuale laboriosa. La forza della regolarizzazione viene invece determinata automaticamente dal vincolo di discrepanza, il quale impone al residuo sui dati di rimanere coerente con il livello di rumore presente nelle misurazioni.

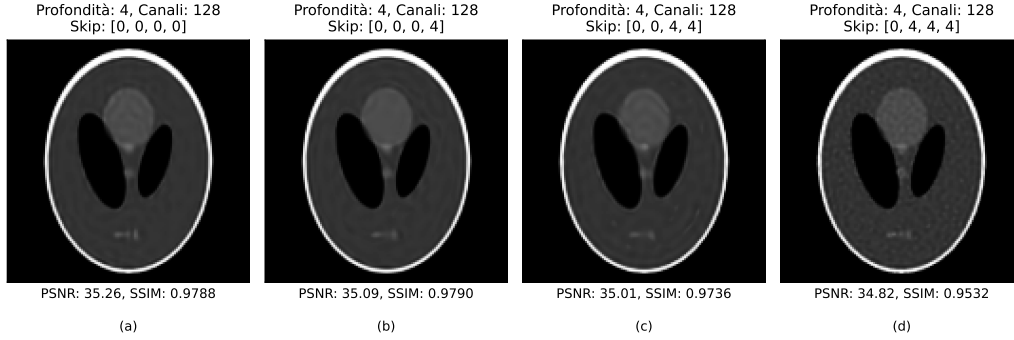


Figura 2.2: Confronto tra le ricostruzioni ottenute con e senza skip connections, per la configurazione con profondità 4 e 128 canali. Per ciascuna variante sono riportati il PSNR e lo SSIM rispetto al phantom di riferimento.

I parametri che rimangono da fissare sono β_t e β_r , i coefficienti di penalizzazione della Lagrangiana aumentata (2.12). A differenza di λ , questi non controllano la forza della regolarizzazione ma influenzano esclusivamente la velocità di convergenza dell'algoritmo. Cascarano et al. [4] mostrano empiricamente che, all'interno di un intervallo ragionevole di valori, la qualità della ricostruzione finale risulta poco sensibile alla loro scelta specifica. Questo comportamento è stato riscontrato anche nell'implementazione sviluppata in questo lavoro, come mostrato in Figura 2.3.

Nell'esperimento sono stati adottati i valori $\beta_t = 1.0$ e $\beta_r = 1.5$, che assegnano un peso leggermente maggiore alla consistenza con i dati misurati rispetto al vincolo sulla variabile ausiliaria t .

Learning rates Come discusso nella formalizzazione del modello, cDIP-RED prevede, oltre all'aggiornamento dei pesi θ , anche due passi α_x e α_y . Il parametro α_x interviene esclusivamente nel calcolo della variabile ausiliaria t^{k+1} , dove scala il contributo del denoiser nel passo prossimale, controllando quindi quanto l'aggiornamento di t si avvicina all'uscita del denoiser rispetto al valore corrente. Il parametro α_y regola invece la velocità di aggiornamento dei moltiplicatori di Lagrange associati ai vincoli del problema, e influenza la velocità con cui il modello converge a scapito della stabilità. Per l'aggiorna-

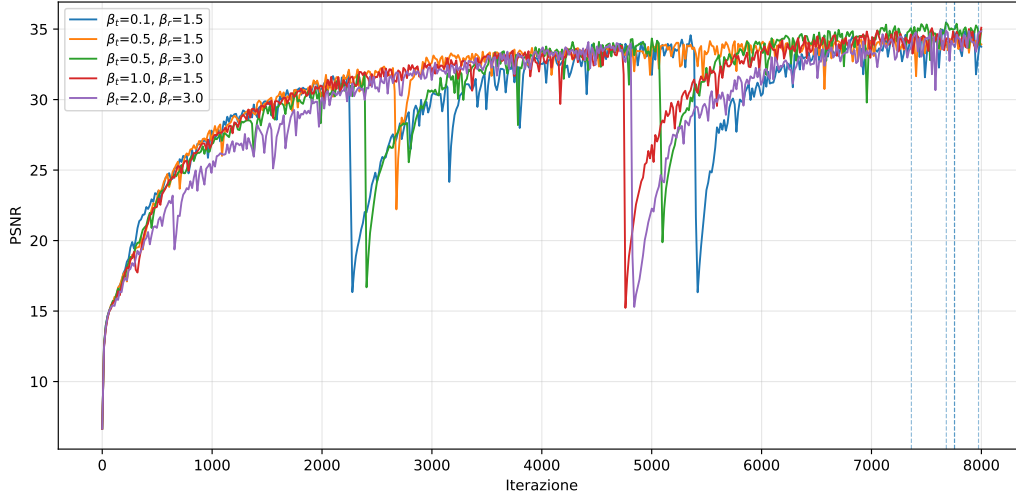


Figura 2.3: Andamento del PSNR al variare dei parametri di penalizzazione β_t e β_r della Lagrangiana aumentata. I valori sono stati campionati all'interno dell'intervallo suggerito in [4].

mento di θ si utilizza l'ottimizzatore Adam [5] con il learning rate di default $\alpha_\theta = 10^{-3}$, in accordo con la configurazione proposta in [4]. I valori adottati per i passi sono $\alpha_x = 5 \times 10^{-1}$ e $\alpha_y = 5 \times 10^{-3}$, in linea con quelli utilizzati dagli autori nell'esperimento di riferimento.

Stima del raggio ρ Come introdotto all'inizio della Sezione 2.3, ρ rappresenta il raggio della palla di vincolo B_ρ e corrisponde alla stima dell'errore presente sulla misura x_0 . Nel caso in cui l'unica sorgente di errore fosse il rumore gaussiano, ρ potrebbe essere stimato analiticamente come $\rho = \sqrt{\tau \cdot m} \cdot \sigma$ [4]. Nel lavoro di questa tesi, tuttavia, l'errore sulla misura è composto da due contributi distinti: il rumore gaussiano aggiunto al sinogramma con deviazione standard $\sigma_\epsilon = 5 \times 10^{-3}$, e l'errore di discretizzazione introdotto dalla differenza tra la geometria ad alta risoluzione usata per costruire x_0 e quella a bassa risoluzione su cui opera la rete. Siccome stimare questi due contributi separatamente sarebbe complesso, ρ viene calcolato direttamente come

$$\rho = \|Ax - x_0\|_2, \quad (2.26)$$

dove x è l'immagine a bassa risoluzione $M \times M$ e A è l'operatore di proiezione corrispondente. In questo modo ρ cattura automaticamente entrambe le sorgenti di errore, fornendo una stima più fedele del reale divario tra la misura e la proiezione dell'immagine ricostruita.

Scelta del denoiser Ricordando il problema (2.11) che il modello cDIP-RED risolve, la scelta del denoiser D è in linea di principio decisiva, poiché interviene direttamente nella funzione obiettivo attraverso la funzione di regolarizzazione RED. In questo lavoro sono stati confrontati il filtro gaussiano e il denoiser basato sulla Variazione Totale (TV). I risultati, riportati in Figura 2.4, mostrano però ricostruzioni di qualità sostanzialmente simile, a indicare che nel modello cDIP-RED la scelta del denoiser pesa meno che in altri metodi, come ad esempio DeepRED [8], dove il denoiser guida più direttamente il risultato. Ciò è probabilmente dovuto al fatto che qui la forza della regolarizzazione è determinata automaticamente dal vincolo di discrepanza, indipendentemente dal denoiser scelto, come del resto rilevato anche da Cascarano et al. [4]. Per questo negli esperimenti del Capitolo 3 si è optato per il filtro gaussiano, che risulta più semplice da un punto di vista computazionale.

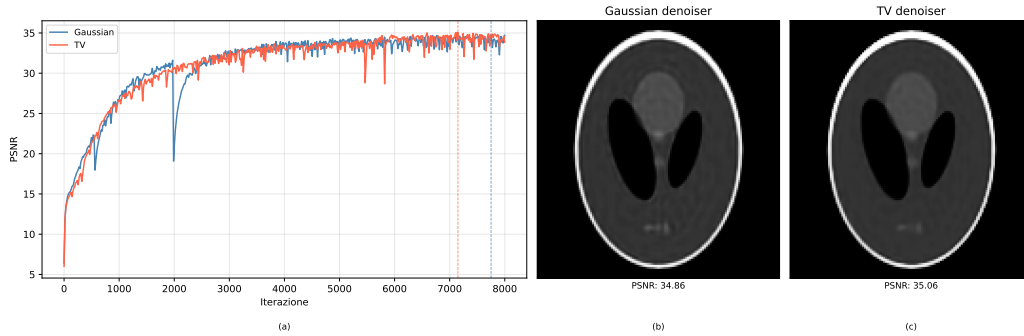


Figura 2.4: Confronto tra denoiser gaussiano e TV nel modello cDIP-RED. Da sinistra: andamento del PSNR al variare delle iterazioni, ricostruzione con denoiser gaussiano, ricostruzione con denoiser TV.

Capitolo 3

Risultati sperimentali

In questo capitolo vengono presentati i risultati ottenuti dall'applicazione del modello cDIP-RED alla ricostruzione tomografica. Vengono prima introdotte le metriche di valutazione utilizzate, per poi discutere il comportamento del metodo durante l'ottimizzazione e le ricostruzioni ottenute al variare del numero di angoli di proiezione disponibili.

3.1 Metriche di valutazione

La metrica principale adottata in questo lavoro è il Peak Signal-to-Noise Ratio (PSNR), che esprime il rapporto tra la potenza del segnale nell'immagine di riferimento e la potenza dell'errore di ricostruzione, ed è definito come:

$$\text{PSNR}(x, \hat{x}) = 10 \log_{10} \frac{x_{\max}^2}{\text{MSE}(x, \hat{x})},$$

dove $x_{\max}$ è il valore massimo dell'immagine di riferimento, e MSE è il Mean Squared Error definito di seguito. Valori del PSNR più alti indicano una ricostruzione più fedele all'originale, e per questo motivo il PSNR viene monitorato durante l'ottimizzazione e impiegato per selezionare la ricostruzione migliore al termine delle iterazioni.

Il Mean Squared Error (MSE), invece, misura l'errore quadratico medio tra la ricostruzione e l'immagine di riferimento. Per un vettore $x \in \mathbb{R}^M$, esso

è definito come

$$\text{MSE}(x, \hat{x}) = \frac{1}{M} \|x - \hat{x}\|_2^2. \quad (3.1)$$

Accanto al PSNR si utilizza lo Structural Similarity Index Measure (SSIM, introdotto nella Sezione 2.4, che valuta la qualità percettiva dell'immagine tenendo conto di luminanza, contrasto e struttura locale. A differenza del PSNR, l'SSIM è sensibile alle correlazioni spaziali tra i pixel, rendendolo più coerente con la percezione visiva umana. Il valore di SSIM è compreso in $[-1, 1]$, con $\text{SSIM} = 1$ che indica identità perfetta tra le due immagini.

Infine, si utilizza il Constraint Ratio (CR) [4], che misura il rapporto tra il residuo sui dati e la stima del rumore, permettendo di verificare il soddisfacimento del vincolo di discrepanza durante l'ottimizzazione. Dato $f_\theta(z)$ l'output della rete al passo corrente, esso è definito come:

$$\text{CR}(f_\theta(z)) = \frac{\|Af_\theta(z) - x_0\|_2^2}{\delta^2}, \quad (3.2)$$

dove $\delta^2 = \sigma^2 N$ è il valore atteso del residuo al quadrato sotto l'ipotesi di rumore gaussiano. Il vincolo di discrepanza risulta soddisfatto quando $\text{CR} \leq 1$.

3.2 Convergenza del modello

Si mostrano ora i risultati della ricostruzione tomografica ottenuta tramite l'implementazione del modello cDIP-RED sviluppata in questo lavoro di tesi. Il problema di test è costruito come descritto dal sistema (2.3) nella Sezione 2.1, calcolando il sinogramma proiettando una versione del phantom di Shepp-Logan a risoluzione doppia rispetto a quella usata nella ricostruzione, al fine di evitare l'inverse crime. La geometria di acquisizione prevede $K = 60$ angoli di proiezione equidistanti nell'intervallo $[0^\circ, 180^\circ]$ e $P = 182$ raggi X per angolo. Al sinogramma risultante è stato applicato rumore gaussiano con deviazione standard pari allo 0.5% del massimo del segnale.

Il modello cDIP-RED viene eseguito per 10000 iterazioni, selezionando come ricostruzione finale l'immagine $x = f_\theta(z)$ corrispondente al massimo

PSNR osservato durante l'ottimizzazione. In Figura 3.1 si confronta la ricostruzione così ottenuta con quella prodotta dal metodo FBP sulla stessa acquisizione.

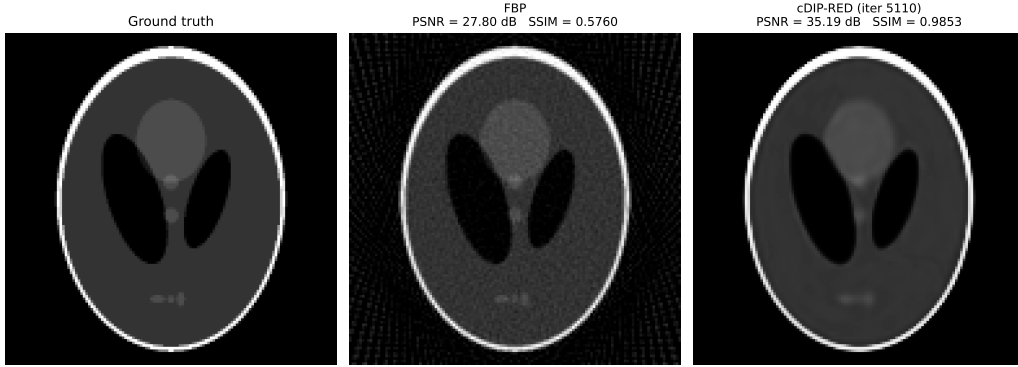


Figura 3.1: Confronto tra le ricostruzioni del phantom di Shepp-Logan.

Le metriche quantitative mostrano un netto vantaggio di cDIP-RED rispetto all'FBP: il PSNR aumenta da 27.80 dB a 35.19 dB, con un guadagno di 7.39 dB, mentre l'SSIM cresce da 0.5760 a 0.9853, avvicinandosi sensibilmente al valore unitario. Tuttavia, a un'analisi visiva diretta, la ricostruzione FBP appare localmente più nitida ai bordi delle strutture, mentre quella prodotta da cDIP-RED risulta leggermente più sfumata, effetto attribuibile all'azione del regolarizzatore RED che penalizza componenti ad alta frequenza nell'immagine ricostruita.

In Figura 3.2 si riportano le tre metriche monitorate durante le 10000 iterazioni del modello.

Il comportamento del PSNR in Figura 3.2 (a) è particolarmente rilevante. A differenza del metodo DIP classico, in cui il PSNR raggiunge un picco per poi degradare rapidamente richiedendo un arresto prematuro dell'ottimizzazione, cDIP-RED non mostra alcun fenomeno di overfitting progressivo: dopo la flessione attorno all'iterazione 5500, la metrica continua a crescere. Questo risultato è in accordo con quanto dimostrato sperimentalmente da Cascarano et al. [4], e conferma che la formulazione vincolata con il principio di discrepanza di Morozov elimina la dipendenza dall'early stopping.

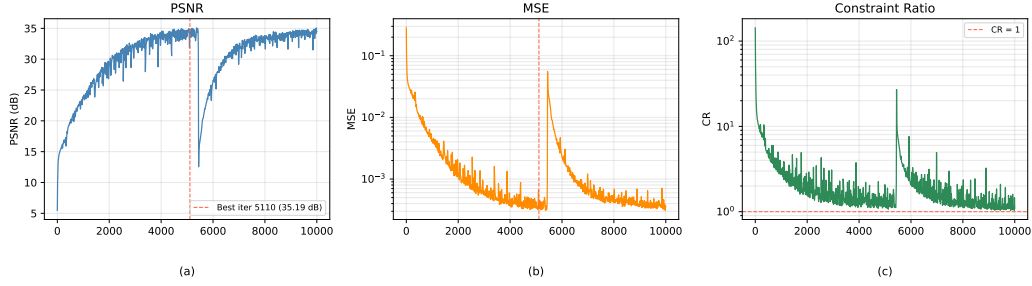


Figura 3.2: Andamento delle metriche di ottimizzazione durante le 10000 iterazioni del modello cDIP-RED.

Per verificare questo comportamento su un orizzonte temporale più esteso, si è eseguito il modello per 20000 iterazioni. I risultati sono riportati in Figura 3.3.

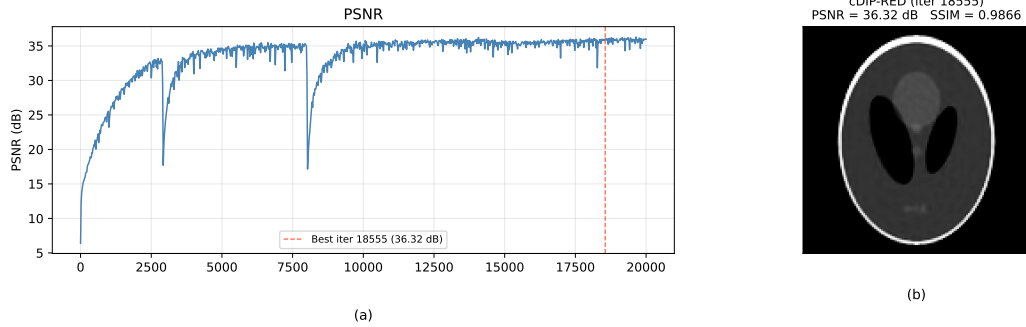


Figura 3.3: Risultati del modello cDIP-RED su 20000 iterazioni.

Estendendo l'ottimizzazione a 20000 iterazioni, il PSNR raggiunge 36.32 dB all'iterazione 18555, con uno SSIM di 0.9866, portando il guadagno totale rispetto all'FBP a 8.52 dB in PSNR e a 0.4106 in SSIM. La curva di PSNR in Figura 3.3 (a) non presenta alcun plateau o inversione di tendenza nelle 20000 iterazioni considerate, confermando che cDIP-RED sfrutta ogni iterazione aggiuntiva in modo produttivo. Infine, guardando la Figura 3.2 (c), si verifica che il modello soddisfa correttamente il vincolo di discrepanza: il Constraint Ratio decresce monotonicamente da valori dell'ordine di 10^2 verso l'unità, indicando che il residuo $\|Af_\theta(z) - x_0\|_2^2$ converge a $\delta^2 N$, in accordo con il

principio di discrepanza di Morozov adottato nella formulazione (2.11).

Conclusioni

La ricostruzione tomografica da un numero ridotto di proiezioni è un problema inverso intrinsecamente mal posto, in quanto piccole perturbazioni nel sinogramma misurato producono soluzioni molto diverse tra loro, rendendo necessario introdurre informazioni aggiuntive sull'immagine cercata attraverso la regolarizzazione. In questo lavoro è stato quindi studiato e implementato il metodo cDIP-RED di Cascarano et al. [4], che affronta il problema combinando il prior implicito dell'architettura di una rete generatrice con la regolarizzazione RED, all'interno di un problema di ottimizzazione vincolata che elimina la necessità di calibrare manualmente il parametro di regolarizzazione λ .

Il punto di partenza è stato il modello di acquisizione tomografica, ricavato formalizzando la Trasformata di Radon e la Legge di Beer-Lambert e arrivando al sistema lineare discreto $Ax = x_0$. Questo sistema è quasi sempre sottodeterminato e mal condizionato quando le proiezioni disponibili sono poche, e la sua inversione diretta, come mostra già il metodo FBP in condizioni di scarso campionamento angolare, produce artefatti da streaking che compromettono la qualità dell'immagine ricostruita.

L'approccio DIP risolve il problema di regolarizzazione in modo non convenzionale: invece di definire esplicitamente una funzione di regolarizzazione $\mathcal{R}(x)$, forza l'immagine ricostruita a essere l'output di una rete generatrice convoluzionale f_θ , la cui architettura codifica un prior implicito sulle immagini naturali. Il limite di questa formulazione è la dipendenza dall'early stopping, in quanto senza un criterio oggettivo di arresto, l'ottimizzazione

finisce per adattarsi al rumore. Il modello cDIP-RED supera questo limite riformulando il problema come ottimizzazione vincolata secondo il principio di discrepanza di Morozov, imponendo che il residuo $\|Af_\theta(z) - x_0\|_2^2$ rimanga coerente con il livello di rumore stimato $\delta^2 N$. In questo modo la forza della regolarizzazione viene determinata automaticamente dal vincolo, e l'unica quantità da stimare diventa δ , il cui calcolo è notevolmente più agevole rispetto alla calibrazione manuale di λ . Il problema minimax risultante viene risolto tramite una versione modificata del metodo PGDA, di cui è stato derivato lo schema iterativo esplicito.

I risultati sperimentali ottenuti mostrano un vantaggio netto di cDIP-RED rispetto all'FBP sia in termini di PSNR che di SSIM. Inoltre, estendendo l'ottimizzazione a un numero elevato di iterazioni la qualità della ricostruzione continua a migliorare senza segni di degradazione. Questo è forse il risultato più interessante dell'intera analisi, ovvero l'andamento crescente del PSNR per l'intera durata dell'ottimizzazione, senza il crollo tipico del DIP classico, dimostra sperimentalmente che il vincolo di discrepanza svolge davvero il ruolo di regolarizzatore stabile, sostituendo a tutti gli effetti l'early stopping. Lo conferma anche il Constraint Ratio, che nel corso delle iterazioni decresce progressivamente verso l'unità, mostrando che il residuo si stabilizza attorno al livello di rumore atteso.

La criticità emersa riguarda la nitidezza visiva della ricostruzione. Nonostante le metriche siano buone, l'immagine prodotta da cDIP-RED appare leggermente più sfumata rispetto a quanto ci si aspetterebbe, probabilmente perché il regolarizzatore gaussiano utilizzato nel denoiser sopprime componenti ad alta frequenza che corrispondono a bordi reali dell'immagine. Un possibile sviluppo futuro consiste nel sostituire il denoiser gaussiano con uno più orientato alla preservazione dei bordi, ad esempio basato sulla Variazione Totale o su un denoiser appreso, verificando se la stabilità garantita dal vincolo di discrepanza si mantiene inalterata.

Bibliografia

- [1] Daniel Otero Baguer, Johannes Leuschner, and Maximilian Schmidt. Computed tomography reconstruction using deep image prior and learned reconstruction methods. *Inverse Problems*, 36(9):094004, September 2020.
- [2] Stephen Boyd and Lieven Vandenberghe. *Convex Optimization*. Cambridge University Press, 2004. Chapter 5: Duality (pp. 215-273).
- [3] Radu Ioan Boţ and Axel Böhm. Alternating proximal-gradient steps for (stochastic) nonconvex-concave minimax problems. *SIAM Journal on Optimization*, 33(3):1884–1913, August 2023.
- [4] Pasquale Cascarano, Giorgia Franchini, Erich Kobler, Federica Porta, and Andrea Sebastiani. Constrained and unconstrained deep image prior optimization models with automatic regularization. *Computational Optimization and Applications*, 84(1):125–149, July 2022.
- [5] Ian Goodfellow, Yoshua Bengio, and Aaron Courville. *Deep Learning*. MIT Press, 2016. <http://www.deeplearningbook.org>.
- [6] Avinash C Kak and Malcolm Slaney. *Principles of computerized tomographic imaging*. Society for Industrial and Applied Mathematics, January 2001.
- [7] Victor Lempitsky, Andrea Vedaldi, and Dmitry Ulyanov. Deep image prior. In *2018 IEEE/CVF Conference on Computer Vision and Pattern Recognition*, page 9446–9454. IEEE, June 2018.

-
- [8] Gary Mataev, Michael Elad, and Peyman Milanfar. Deepred: Deep image prior powered by red, 2019.
 - [9] Jennifer L Mueller and Samuli Siltanen. *Linear and nonlinear inverse problems with practical applications*. SIAM, 2012.
 - [10] Chiyuan Zhang, Samy Bengio, Moritz Hardt, Benjamin Recht, and Oriol Vinyals. Understanding deep learning requires rethinking generalization, 2016.