

UNIVERSITA' DEGLI STUDI DI PADOVA

FACOLTA' DI SCIENZE MM.FF.NN.

Corso di Laurea in Matematica
Dipartimento di Matematica Pura ed Applicata

TESI DI LAUREA

**ANALISI ED IMPLEMENTAZIONE DI
DI UN SOLUTORE NUMERICO PER EQUAZIONI
INTEGRALI NON LINEARI DELLA
TEORIA DEL TRASPORTO**

Relatore: Dott. Marco Vianello

Correlatore: Dott. Alvise Sommariva

Laureando: Enrico Facchinello

ANNO ACCADEMICO 1997-98

Capitolo 1

Preliminari

1.1 Una classe di equazioni di Hammerstein della teoria del trasporto.

Equazioni integrali non lineari quadratiche del tipo

$$a(x)\phi(x) + \phi(x) \int_{\Omega} \mathcal{K}(x, t)\phi(t) d\mu(t) = b(x), \quad x \in \Omega, \quad (1.1)$$

dove (Ω, S, μ) è uno spazio di misura, si presentano in modo naturale nel contesto della teoria del trasporto. In (1.1), il kernel $\mathcal{K}$, e i coefficienti a, b , sono funzioni $(\mu-)$ misurabili e non negative $(\mu-q.o.)$. Da un punto di vista fisico, è interessante cercare soluzioni positive $(\mu-)$ integrabili (soluzioni $L_+^1(\Omega, \mu)$), eventualmente continue quando Ω è dotato di una topologia di Hausdorff, visto che ϕ ha generalmente il significato di una funzione di distribuzione. Si può notare che la formulazione generale (1.1) ci permette di trattare contemporaneamente sia i casi in dimensione infinita che le corrispondenti versioni finito-dimensionale, prodotte, ad esempio, dalla discretizzazione per scopi numerici. In questo caso, possiamo pensare o che μ sia una misura discreta sul continuo $\Omega \subseteq \mathbb{R}^d$, o più semplicemente che Ω sia uguale a $\{1, \dots, m\}$, μ sia la misura che conta i punti, a, b , e ϕ siano vettori m -dimensionali e che K sia una matrice $m \times m$.

È noto che importanti modelli nella teoria del trasferimento radiativo e del trasporto di neutroni [5], come anche nella teoria cinetica dei gas [9, 10],

sono descritti da equazioni integrali non lineari della forma:

$$h(x) = y(x) + h(x) \int_{\Omega} \mathcal{K}(x, t) h(t) d\mu(t), \quad x \in \Omega, \quad (1.2)$$

dove rispettivamente $\Omega = [0, 1]$ oppure $\Omega = (0, \infty)$, $y \in L_+^1(D)$, e $0 \leq \mathcal{K}(x, t) \leq 1$, $\mathcal{K}(x, t) + \mathcal{K}(t, x) \equiv 1$. Quando $\|y\|_1 = \int_{\Omega} y(t) d\mu(t) < 1/2$, ipotesi che corrisponde ai casi non conservativi, non è difficile vedere che l'equazione (1.2) può essere trasformata in

$$(1 - 2\|y\|_1)^{1/2} h(x) + h(x) \int_{\Omega} \mathcal{K}(t, x) h(t) d\mu(t) = y(x), \quad x \in \Omega, \quad (1.3)$$

che chiaramente appartiene alla stessa famiglia di (1.1). Si noti che l'operatore a secondo membro in (1.2) è (puntualmente) crescente su argomenti h non negativi, mentre se si scrive (1.3) come una equazione di punto fisso, si ottiene un operatore decrescente. All'interno della famiglia di equazioni del tipo (1.2) e (1.3), c'è per esempio la famosa equazione di Chandrasekhar [5], che ha un ruolo chiave nella modellizzazione del trasferimento radiativo nelle atmosfere semi-infinite, dove $D = [0, 1]$, $\mathcal{K}(x, t) = x/(x + t)$, e y è costante.

Nell'ambito della formulazione cosiddetta "scattering-kernel" per l'equazione di Boltzmann, che fornisce un modello per la diffusione di particelle in una miscela di due specie (particelle di campo e particelle test), lo studio del caso omogeneo nello spazio in assenza di forze porta in modo naturale all'equazione integrale quadratica

$$N\hat{g}_r(|\mathbf{u}|)f(\mathbf{v}) + f(\mathbf{v}) \int_{\mathbb{R}^3} g_r(|\mathbf{v} - \mathbf{v}'|)f(\mathbf{v}')d\mathbf{v}' = Q_0 S(\mathbf{v}), \quad \mathbf{v} \in \mathbb{R}^3, \quad (1.4)$$

che corrisponde alla ricerca di soluzioni stazionarie. In (1.4), $\hat{g}$ e g descrivono l'effetto di rimozione tra le particelle, N è la densità totale fissa delle particelle di campo e $Q_0 S(\mathbf{v})$ rappresenta la velocità di emissione delle particelle test (assunta costante nel tempo e uniforme nello spazio), con intensità Q_0 , dove $\int_{\mathbb{R}^3} S(\mathbf{v}) d\mathbf{v} = 1$. Nel caso isotropo, cioè $S(\mathbf{v}) = (4\pi)^{-1} S_0(v)$, si può vedere che $f(v) = (4\pi)^{-1} S_0(v)$ risolve (1.4), se $f_0(v)$ è una soluzione di

$$N\hat{g}_r(v)f(v) + f(v) \int_0^\infty G(v, v')f(v') dv' = Q_0 S_0(v), \quad v \in (0, \infty), \quad (1.5)$$

dove il kernel G è dato da

$$G(v, v') = \frac{1}{2vv'} \int_{v-v'}^{v+v'} tg_r d\mu(t). \quad (1.6)$$

Si noti che tutte le funzioni e i parametri nei modelli (1.4) e (1.5) sono positivi, o perlomeno non negativi, visto il loro significato fisico. Per la stessa ragione si è interessati ad una soluzione non negativa in L^1 , che rappresenta la distribuzione delle particelle test in funzione della velocità. In particolare, la norma L^1 di f fornisce la densità incognita delle particelle test.

Sotto le ipotesi naturali

$$ab > 0 \text{ q.o. rispetto a } \mu \text{ in } \Omega, \text{ e } b/a, \mathcal{K}(x, \cdot) b/a \in L^1(\Omega, \mu), \quad (1.7)$$

attraverso la trasformazione

$$\phi = \frac{b}{a} \frac{1}{1+u}, \quad (1.8)$$

si riduce l'originaria equazione quadratica alla forma di Hammerstein

$$u(x) = \int_{\Omega} \frac{\mathcal{K}(x, t)b(t)}{a(x)a(t)} \frac{1}{1+u(t)} d\mu(t), \quad x \in \Omega. \quad (1.9)$$

L'ipotesi base $b\mathcal{K}(x, \cdot)/a \in L^1(\Omega, \mu)$, garantisce che l'equazione di Hammerstein (1.9) sia ben definita. Richiedere la positività q.o. rispetto a μ di a e b non è restrittivo, essendo a, b non negative q.o., a meno che non si annullino contemporaneamente su un sottoinsieme di misura non nulla: è sufficiente infatti cambiare il dominio con l'insieme $\{x \in \Omega : a(x)b(x) > 0\}$, poiché $\phi = 0$ dove $b = 0$. Si osservi che ogni soluzione non negativa dell'equazione di Hammerstein (1.9) corrisponde attraverso (1.8) ad una soluzione in $L^1_+(\Omega, \mu)$ dell'equazione quadratica (1.1) (con $\phi(x) \leq b(x)/a(x)$); d'altra parte, se $\phi(x) \geq 0$ risolve (1.1) allora necessariamente $\phi(x) \leq b(x)/a(x)$ (dove tutte le disuguaglianze sono intese q.o. rispetto a μ).

L'equazione (1.9) è un caso particolare dell'equazione di Hammerstein

$$u(x) = A(u)(x) = \int_{\Omega} k(x, t)g(u(t))d\mu(t) + \psi(x) \quad (1.10)$$

dove $k(x, \cdot) \in L^1_+(\Omega, \mu)$ e $g : [0, \infty] \rightarrow [0, \infty]$ è continua e strettamente decrescente con $y \mapsto yg(y)$ strettamente crescente; assumiamo inoltre che $\psi \in L^1_+(\Omega, \mu)$. Le ipotesi su g sono equivalenti a

$$g(y) = \frac{1}{c + q(y)}, \quad c > 0, \quad (1.11)$$

dove $q \geq 0$ ha le seguenti proprietà: è strettamente crescente in $[0, \infty)$, esiste $\lim_{y \rightarrow +\infty} q(y)/y > 0$, ed è strettamente *sublineare*, cioè, $q(\tau y) > \tau q(y)$ per ogni $\tau \in (0, 1)$, $y > 0$. Un esempio importante di tale q è dato da $q(y) = \sum_{i=1}^m c_i y^{\alpha_i}$, dove $c_i > 0$, $0 < \alpha_i < \dots < \alpha_m \leq 1$.

Studieremo l'equazione (1.10) in due situazioni distinte:
nel primo caso supporremo che esista β μ -misurabile positiva tale che

$$k(x, t)\beta(t) = k(t, x)\beta(x) \text{ q.o. in } \Omega \times \Omega. \quad (1.12)$$

Nel secondo caso si considererà l'ipotesi

$$0 < \sigma := \text{ess sup}_{\Omega} \int_{\Omega} k(x, t) d\mu(t) < \infty. \quad (1.13)$$

Fermiamoci prima un attimo per introdurre alcuni concetti che saranno utili in seguito.

1.2 Spazi di Banach ordinati

1.2.1 Coni

Sia E uno spazio di Banach. Un insieme chiuso non vuoto $P \subseteq E$ è detto “cono” se

- $x \in P$, $\lambda \geq 0$ implica che $\lambda x \in P$;
- $x \in P$, $-x \in P$, implica che $x = \theta$ (qui θ indica l'elemento nullo di E).

È importante notare che un cono definisce un ordinamento parziale in E definendo

$$x \preceq y \Leftrightarrow y - x \in P;$$

per questa ragione la coppia (E, P) è chiamata *spazio di Banach ordinato*. Data questa definizione, se $x \preceq y$ e $x \neq y$ scriviamo $x \prec y$, e con $[a, b]$ intendiamo l'intervallo contenente tutti gli x in E tali che $a \preceq x \preceq b$.

Diciamo che un cono è

- *solido* se contiene punti interni;

- *normale* se per tutte le successioni $\{x_n\}, \{y_n\}, \{z_n\}$ in P tali che $x_n \preceq z_n \preceq y_n$, $\lim_{n \rightarrow +\infty} \|x_n - x\| = 0$, $\lim_{n \rightarrow +\infty} \|y_n - x\| = 0$ (per un certo $x \in P$) abbiamo $\lim_{n \rightarrow +\infty} \|z_n - x\| = 0$ (cioè vale *il teorema dei due carabinieri*); una condizione necessaria e sufficiente perchè il cono P sia normale è che esista una norma equivalente $\|\cdot\|_1$ (rispetto a $\|\cdot\|_E$) tale che $\theta \preceq x \preceq y$ implichi $\|x\|_1 \preceq \|y\|_1$;
- *regolare* se per ogni successione $\{x_n\}$ in E limitata rispetto all'ordine, cioè tale che

$$x_1 \preceq \dots \preceq x_n \preceq \dots \preceq y, \quad n = 1, 2, \dots \quad y \in E$$

esiste x in E tale che $\lim_{n \rightarrow +\infty} \|x_n - x\| = 0$.

Ci sono interessanti relazioni tra i coni normali e i coni regolari (vedi [8]).

Proposizione 1.2.1 *Sia E uno spazio di Banach e P un cono in E . Se P è regolare allora P è normale. Il viceversa è vero se E è riflessivo.*

Ora consideriamo qualche esempio di coni in tre noti spazi di Banach ordinati, vedendo se sono solidi, normali o regolari.

Esempio 1.2.1 *Sia $E = \mathbb{R}^m$ munito della norma euclidea, e $P = \{x = (x_1, \dots, x_m) : x_i \geq 0, i = 1 \dots m\}$. Il cono P è solido e normale (poiché la norma è monotona). Inoltre, essendo $\mathbb{R}^m$ riflessivo, P è anche regolare.*

Esempio 1.2.2 *Sia $E = C[0, 1]$ lo spazio delle funzioni continue nell'intervallo $[0, 1]$, munito della "sup-norma" e $P = C_+[0, 1] = \{h \in C[0, 1] : h(x) \geq 0 \text{ per ogni } x \in [0, 1]\}$. Il cono P è solido e normale, ma non regolare. Infatti, sia $h_n(x) = 1 - x^n$, $x \in [0, 1]$, ($n = 1, 2, 3, \dots$). Chiaramente*

$$h_1(x) \preceq h_2(x) \preceq \dots \preceq h_n(x) \preceq \dots \preceq 1$$

ma $\{h_n\}$ non converge in $C[0, 1]$.

Esempio 1.2.3 *Sia $E = L^p(\Omega)$ lo spazio delle funzioni Lebesgue misurabili con potenza p -esima sommabile su $\Omega \subset \mathbb{R}^m$, dove $p \geq 1$ intero finito, $0 < \text{mis}(\Omega) < \infty$. Il cono $P = L_+^p(\Omega) = \{h \in L^p(\Omega) : h(x) \geq 0 \text{ per ogni } x \in \Omega\}$ è regolare ma non solido; per la Proposizione 1.2.1, P è normale.*

1.2.2 Operatori monotoni

Sia (E, P) uno spazio di Banach ordinato, con $\preceq$ ordine parziale definito dal cono P .

Un operatore $A : D \subseteq E \rightarrow E$ è chiamato

- *sublineare* se $A(\tau u) \succeq \tau A(u)$ per ogni $\tau \in (0, 1)$, $u \succ 0$;
- *crescente* se $u_1, u_2 \in D$, $u_1 \preceq u_2$ implica che $A(u_1) \preceq A(u_2)$;
- *decescente* se $u_1, u_2 \in D$, $u_1 \preceq u_2$ implica che $A(u_2) \preceq A(u_1)$.

Diciamo che $u, v \in D$ sono *punti fissi* di A se $u = A(v)$, $v = A(u)$.

Un operatore $M : D \times D \subseteq E \rightarrow E$ è *mixed monotone* se $M(u, v)$ è crescente in u e decrescente in v , cioè

- (i) $u_1, u_2 \in D$, $u_1 \preceq u_2$ implica che $M(u_1, v) \preceq M(u_2, v)$ per ogni $v \in D$;
- (ii) $v_1, v_2 \in D$, $v_1 \preceq v_2$ implica che $M(u, v_2) \preceq M(u, v_1)$ per ogni $u \in D$.

Ogni $u, v \in D \times D$ tali che $u = M(u, v)$, $v = M(v, u)$ sono chiamati *punti fissi quasi accoppiati* di M .

Capitolo 2

Convergenza dei metodi rilassati di tipo Picard in $\mathcal{M}_+$ e L_+^∞ .

In questo capitolo mostreremo che per u_0 appartenente allo spazio delle funzioni misurabili positive $\mathcal{M}_+(\Omega, \mu)$, la successione $\{u_n\}$ definita da

$$\begin{cases} u_{n+1/2} = E(u_{n+1/2}) + F(u_n), \\ u_{n+1} = (1 - \omega)u_n + \omega u_{n+1/2}, \end{cases} \quad (2.1)$$

(dove $\omega \in (0, 1]$, E ed F sono opportuni operatori decrescenti tali che $A = E + F$), converge all'unica soluzione dell'equazione di Hammerstein (1.10) in $\mathcal{M}_+(\Omega, \mu)$ (ed eventualmente in $L_+^\infty(\Omega, \mu)$).

2.1 Convergenza a soluzioni positive misurabili.

Sia $\mathcal{M}(\Omega, \mu)$ lo spazio delle funzioni misurabili su $\Omega \subseteq \mathbb{R}^d$ modulo la relazione di equivalenza quasi ovunque. $\mathcal{M}(\Omega, \mu)$ può essere ordinato parzialmente dal cono chiuso $\mathcal{M}_+(\Omega, \mu) = \{u \in \mathcal{M}(\Omega, \mu) : u(x) \geq 0 \text{ q.o.}\}$, cioè $v \succeq u$ se $v - u \in \mathcal{M}_+(\Omega, \mu)$. Questo cono è regolare, cioè ogni successione monotona e limitata rispetto a questo ordine in $\mathcal{M}_+(\Omega, \mu)$ ha limite, ed è pure normale, nel senso che vale il teorema dei due Carabinieri. Consideriamo l'equazione

di Hammerstein

$$u(x) = \int_{\Omega} k(x, t) g(u(t)) d\mu(t) + \psi(x), \quad x \in \Omega \quad (2.2)$$

e sia $A : \mathcal{M}_+(\Omega, \mu) \rightarrow \mathcal{M}_+(\Omega, \mu)$ l'operatore definito dal secondo membro di (2.2) dove supporremo che

- (I) Il nucleo $k \in L_+^1(\Omega \times \Omega, \mu \times \mu)$, ($\mu \times \mu$ è la misura prodotto definita da μ) cioè $h \in L^1(\Omega \times \Omega, \mu \times \mu)$ e $k(x) \geq 0$ quasi ovunque;
- (II) $g : \mathbb{R}^+ \rightarrow \mathbb{R}^+$ è una funzione decrescente, tale che $yg(y)$ sia strettamente decrescente q.o.;
- (III) $\psi \in L_+^1(\Omega, \mu)$, ovvero $\psi \in L^1(\Omega, \mu)$, $\psi(x) \geq 0$ q.o. e $\beta\psi \in L_+^1(\Omega, \mu)$;
- (IV) $\beta \left(g(0) \int_{\Omega} k(\cdot, t) d\mu(t) + \psi \right) g \left[g(0) \int_{\Omega} k(\cdot, t) d\mu(t) + \psi \right] \in L^1(\Omega, \mu)$.

Con queste assunzioni è immediato notare che l'operatore A è decrescente, cioè $u \preceq v$ implica che $A(u) \succeq A(v)$. In questo paragrafo supporremo inoltre che l'operatore A verifichi la proprietà (1.12).

Cercheremo una soluzione non negativa di (2.2) in qualsiasi sottospazio di $\mathcal{M}_+(\Omega, \mu)$. In simili contesti, usualmente si assume l'ipotesi (1.13) che viene ora trascurata e sostituita con (1.12), in cui non è richiesto che l'estremo superiore sia infinito.

Teorema 2.1.1 *Supponiamo che siano verificate (I), (II), (III), (IV) e che $M : \mathcal{M}_+(\Omega, \mu) \times \mathcal{M}_+(\Omega, \mu) \rightarrow \mathcal{M}_+(\Omega, \mu)$ sia l'operatore*

$$M(u, v) = (1 - \omega)u + \omega A(v), \quad (2.3)$$

dove $u, v \in \mathcal{M}_+(\Omega, \mu)$ e $\omega \in (0, 1]$. M ha un unico punto fisso $u^* \in \mathcal{M}_+(\Omega, \mu)$ (cioè $M(u^*, u^*) = u^*$) e, fissati $s_0, t_0 \in [\theta, A(\theta)]$, le successioni $\{s_n\}$, $\{t_n\}$ definite da

$$s_{n+1} = M(s_n, t_n), \quad t_{n+1} = M(t_n, s_n), \quad n = 0, 1, 2, \dots, \quad (2.4)$$

convergono a u^* puntualmente q.o.. In particolare, quando $s_0 = \theta$ e $t_0 = A(\theta)$, vale la seguente catena di disuguaglianze:

$$\theta = s_0 \preceq s_1 \preceq \dots \preceq s_n \preceq \dots \preceq u^* \preceq \dots \preceq t_n \preceq \dots t_1 \preceq t_0 = A(\theta). \quad (2.5)$$

Dimostrazione. Proviamo innanzitutto che

$$\theta = s_0 \preceq s_1 \preceq \dots \preceq s_n \preceq \dots \preceq t_n \preceq \dots \preceq t_1 \preceq t_0 = A(\theta). \quad (2.6)$$

Sappiamo che A è decrescente, dunque se $u, v \in [\theta, A(\theta)]$ si ha che $\theta \preceq A(u) \preceq A(\theta)$ e $\theta \preceq M(u, v) \preceq A(\theta)$. Mostriamo (2.6) per induzione. Per $n = 1$ basta notare che

$$\theta = s_0 \preceq M(s_0, t_0) = s_1 \preceq M(t_0, t_0) \preceq M(s_0, t_0) = t_1 \preceq t_0 = A(\theta).$$

Supponiamo ora che per ipotesi induttiva sia

$$s_{n-1} \preceq s_n \preceq t_n \preceq t_{n-1}$$

Allora

$$s_n = M(s_{n-1}, t_{n-1}) \preceq M(s_n, t_{n-1}) \preceq M(s_n, t_n) = s_{n+1}$$

$$t_n = M(t_{n-1}, s_{n-1}) \succeq M(t_n, s_{n-1}) \succeq M(t_n, s_n) = t_{n+1}$$

$$s_{n+1} = M(s_n, t_n) \preceq M(t_n, t_n) \preceq M(t_n, s_n) = t_{n+1}$$

cioè (2.6) è soddisfatta. Osserviamo che essendo $\mathcal{M}_+(\Omega, \mu)$ regolare, $\{s_n\}$ crescente e limitata da t_0 , $\{t_n\}$ è decrescente e limitata da s_0 , esistono $u^*, v^* \in \mathcal{M}_+(\Omega, \mu)$ tali che $s_n \uparrow u^*$ e $t_n \downarrow v^*$.

Abbiamo dunque che per qualsiasi $n \in \mathbb{N}$

$$\theta \preceq s_1 \preceq \dots \preceq s_n \preceq \dots \preceq u^* \preceq v^* \preceq \dots \preceq t_n \preceq \dots \preceq t_1 \preceq A(\theta), \quad (2.7)$$

e

$$\theta \preceq \lim_{n \rightarrow \infty} s_n = u^* \preceq v^* = \lim_{n \rightarrow \infty} t_n \preceq A(\theta) \text{ puntualmente.}$$

Mostriamo che $u^* = v^*$. Da (2.4)

$$s_{n+1} = M(s_n, t_n) = (1 - \omega)s_n + \omega \int_{\Omega} k(\cdot, t)g(t_n(t)) d\mu(t) + \omega \psi \quad (2.8)$$

e poiché g è strettamente decrescente e $k \in L_+^1(\Omega \times \Omega, \mu \times \mu)$

$$\left| k(x, \cdot)g(t_n(\cdot)) \right| \leq |k(x, \cdot)g(0)| \in L_+^1(\Omega, \mu);$$

dunque essendo g continua

$$\lim_{n \rightarrow \infty} k(x, t)g(t_n(t)) = k(x, t)g(v^*(t)) \text{ perquasiogni } (x, t) \in \Omega \times \Omega$$

e per il teorema della convergenza dominata

$$\lim_{n \rightarrow \infty} \int_{\Omega} k(\cdot, t) g(t_n(t)) d\mu(t) + \psi = \int_{\Omega} k(\cdot, t) g(v^*) d\mu(t) + \psi = A(v^*). \quad (2.9)$$

Visto che $s_n \uparrow u^*$, $t_n \downarrow v^*$, da (2.8) e (2.9)

$$u^* = (1 - \omega)u^* + \omega A(v^*)$$

cioè

$$u^* = A(v^*).$$

Allo stesso modo si vede che $v^* = A(u^*)$. Vogliamo mostrare che in realtà $u^* = v^*$.

Da

$$\theta \preceq u^* \preceq v^* \preceq A(\theta) = g(0) \int_{\Omega} k(\cdot, t) d\mu(t) + \psi$$

essendo $yg(y)$ strettamente crescente, per (2.7) e (IV)

$$\beta u^* g(u^*) \preceq \beta A(\theta) g(A(\theta)) \in L^1(\Omega, \mu). \quad (2.10)$$

Siccome $u^* = A(v^*)$, $v^* = A(u^*)$, necessariamente

$$\beta(x) u^*(x) g(u^*(x)) =$$

$$\beta(x) g(u^*(x)) \int_{\Omega} k(x, t) g(v^*(t)) d\mu(t) + \beta(x) g(u^*(x)) \psi(x).$$

e

$$\beta(x) v^*(x) g(v^*(x)) =$$

$$\beta(x) g(v^*(x)) \int_{\Omega} k(x, t) g(u^*(t)) d\mu(t) + \beta(x) g(v^*(x)) \psi(x) :$$

visto che i termini del primo membro di entrambe le equazioni sono integrabili per (2.10), possiamo sottrarre membro a membro ed integrare su Ω rispetto a $\mu(x)$. Applicando il teorema di Fubini, (ricordando (III)), si vede subito che

$$\begin{aligned} & \int_{\Omega} \beta(x) [u^*(x) g(u^*(x)) - v^*(x) g(v^*(x))] d\mu(x) = \\ & = \int_{\Omega \times \Omega} \beta(x) \left[g(u^*(x)) k(x, t) g(v^*(t)) - g(v^*(x)) k(x, t) g(u^*(t)) \right] d\mu(x) d\mu(t) \end{aligned}$$

$$+ \int_{\Omega} \beta(x) \psi(x) [g(u^*(x)) - g(v^*(x))] d\mu(x). \quad (2.11)$$

Ricordando (1.12)

$$\begin{aligned} & \int_{\Omega \times \Omega} \beta(x) g(u^*(x)) k(x, t) g(v^*(t)) d\mu(x) d\mu(t) = \\ & \int_{\Omega \times \Omega} \beta(t) g(u^*(x)) k(t, x) g(v^*(t)) d\mu(x) d\mu(t) \end{aligned}$$

Scambiando t con x (possiamo farlo perchè stiamo integrando su $\Omega \times \Omega$), quest'ultimo integrale è uguale a

$$\int_{\Omega \times \Omega} \beta(x) g(u^*(t)) k(x, t) g(v^*(x)) d\mu(x) d\mu(t)$$

e quindi essendo il primo termine a secondo membro di (2.11) nullo, abbiamo

$$\begin{aligned} & \int_{\Omega} \beta(x) [u^*(x) g(u^*(x)) - v^*(x) g(v^*(x))] d\mu(x) = \\ & \int_{\Omega} \beta(x) \psi(x) [g(u^*(x)) - g(v^*(x))] d\mu(x). \end{aligned} \quad (2.12)$$

Ora supponiamo per assurdo che $u^* \prec v^*$, cioè che $u^*(x) \leq v^*(x)$ su Ω e $u^*(x) < v^*(x)$ su un certo sottinsieme Ω^+ non trascurabile di Ω . Poiché g è decrescente, deduciamo che $g(u^*) - g(v^*) \geq 0$ q.o.. Ne consegue che il secondo membro di (2.12) è non negativo. D'altro canto, essendo $yg(y)$ strettamente decrescente q.o.

$$u^*(x) g(u^*(x)) - v^*(x) g(v^*(x)) \leq v^*(x) g(v^*(x)) - v^*(x) g(v^*(x)) = 0 \quad (2.13)$$

ed in Ω^+ (2.13) vale strettamente, per cui essendo β strettamente positiva, il primo membro di (2.12) è negativo. Assurdo. Quindi $u^* = v^*$.

Vediamo ora l'unicità del punto fisso di M (e quindi di A). Sia $\xi \in \mathcal{M}_+(\Omega, \mu)$ tale che $\xi = A(\xi)$. Essendo

$$M(\xi, \xi) = (1 - \omega)\xi + \omega A(\xi) = (1 - \omega)\xi + \omega\xi = \xi$$

ξ è punto fisso pure di M . Mostriamo che

$$\theta = s_0 \preceq s_1 \preceq \dots \preceq s_n \preceq \xi = A(\xi) \preceq t_n \preceq \dots \preceq t_1 \preceq t_0 = A(\theta) \quad (2.14)$$

per ogni $n \in \mathbb{N}$.

Essendo $A : \mathcal{M}_+(\cdot, \mu) \rightarrow \mathcal{M}_+(\cdot, \mu)$ un operatore decrescente, abbiamo $\theta = s_0 \preceq \xi \preceq t_0 = A(\theta)$. e quindi il primo passo dell'induzione è verificato. Supponiamo ora che (2.14) valga per un certo n ; essendo

$$s_{n+1} = M(s_n, t_n) \preceq M(\xi, t_n) \preceq M(\xi, \xi) = \xi \preceq M(t_n, s_n) = t_{n+1},$$

ricordando (2.6) la tesi induttiva è verificata.

La normalità del cono $\mathcal{M}_+(\Omega, \mu)$ e la convergenza puntuale q.o. di s_n, t_n a u^* , ci permettono di concludere che $\xi = u^*$. $\square$

Teorema 2.1.2 *Supponiamo $E, F : \mathcal{M}_+(\Omega, \mu) \rightarrow \mathcal{M}_+(\Omega, \mu)$ siano operatori decrescenti tali che $A = E + F$. Si assuma che l'equazione di punto fisso $u = E(u) + \eta$ abbia un'unica soluzione in $\mathcal{M}_+(\Omega, \mu)$ per ogni $\eta \in [\theta, F(\theta)] \subseteq [\theta, A(\theta)]$. Allora per $\omega \in (0, 1]$, la successione $\{u_n\}$ definita da*

$$\begin{cases} u_0 \in [\theta, A(\theta)] \\ u_{n+1/2} = E(u_{n+1/2}) + F(u_n), \\ u_{n+1} = (1 - \omega)u_n + \omega u_{n+1/2} \end{cases} \quad (2.15)$$

converge all'unica soluzione u^* di A in $\mathcal{M}_+(\Omega, \mu)$.

Dimostrazione. Fissato $\omega \in (0, 1]$, consideriamo le successioni $\{s_n\}, \{t_n\}$ definite da

$$\begin{cases} t_0 = A(\theta) \\ t_{n+1/2} = E(u_n) + F(u_n) = A(u_n) \\ t_{n+1} = (1 - \omega)s_n + \omega t_{n+1/2}, \end{cases} \quad (2.16)$$

$$\begin{cases} s_0 = \theta \\ s_{n+1/2} = E(u_n) + F(u_n) = A(u_n) \\ s_{n+1} = (1 - \omega)t_n + \omega s_{n+1/2}. \end{cases} \quad (2.17)$$

Se $M : \mathcal{M}_+(\Omega, \mu) \rightarrow \mathcal{M}_+(\Omega, \mu)$ è l'operatore definito da $M(u, v) = (1 - \omega)u + \omega A(v)$, allora le successioni (2.16), (2.17) possono essere riscritte come

$$s_{n+1} = M(s_n, t_n), \quad t_{n+1} = M(t_n, s_n), \quad n = 0, 1, 2, \dots$$

Per il Teorema 2.1.1 convergono a u^* , che è l'unico punto fisso di M (e di A per (2.3)) in $\mathcal{M}_+(\Omega, \mu)$.

Proviamo per induzione che

$$s_n \preceq t_{n+1/2}, \quad s_{n+1/2} \preceq t_{n+1/2}, \quad n = 0, 1, 2, \dots \quad (2.18)$$

Per quanto riguarda la prima disuguaglianza, per (2.16), (2.17), otteniamo $s_0 \preceq t_{1/2}$; assumendo che $s_{n-1} \preceq t_{n-1/2}$, abbiamo per convessità, $s_n = (1 - \omega)s_{n-1} + \omega t_{n-1/2} \preceq t_{n-1/2}$. Essendo $t_n \preceq t_{n-1}$ per (2.5) e A decrescente, da (2.16) otteniamo $t_{n-1/2} = A(t_{n-1/2}) \preceq A(t_{n-1}) \preceq A(t_n) = t_{n+1/2}$, e conseguentemente $s_n \preceq t_{n+1/2}$.

Per la seconda disuguaglianza in (2.18), per (2.16), (2.17) è facile vedere che $t_1 = s_{1/2} = t_0$. Assumendo $s_{n-1/2} \preceq t_n$, da $s_{n+1/2} = A(s_n) \preceq A(s_{n-1} = s_{n-1/2} \preceq t_n)$, otteniamo, ancora per convessità, $s_{n+1/2} \preceq (1 - \omega)t_n + \omega s_{n+1/2} = t_{n+1}$, che conclude il passo induttivo.

Ora proviamo che $\{u_n\}$ converge a u^* . Si noti che questa successione è ben definita, poiché $u = E(u) + b$ ha un'unica soluzione in $\mathcal{M}_+(\Omega, \mu)$, per ogni $b \in [\theta, F(\theta)]$. Per provare per induzione che

$$\theta = s_0 \preceq s_1 \preceq \dots \preceq s_n \preceq u_n \preceq t_n \preceq \dots \preceq t_1 \preceq t_0 = A(\theta); \quad (2.19)$$

iniziamo mostrando che

$$s_0 \preceq s_1 \preceq u_1 \preceq t_1 \preceq t_0. \quad (2.20)$$

Per il fatto che $u_{1/2} \succeq s_0$, $u_0 \succeq s_0$ otteniamo per (2.15) e (2.17) che $u_{1/2} \preceq s_{1/2} = A(\theta) = t_0$, $u_0 \preceq t_0$, che insieme a (2.15) e (2.16) implica $t_{1/2} \preceq u_{1/2}$. Ora $t_{1/2} \preceq u_{1/2} \preceq s_{1/2}$, $s_0 \preceq u_0 \preceq t_0$, e ancora (2.15), (2.17), (2.16) danno per convessità (2.20).

Assumiamo (2.19) come ipotesi induttiva: prima proviamo che $u_{n+1} \preceq t_{n+1}$. Per (2.15) e (2.17), questo è vero se $u_n \preceq t_n$ e $u_{n+1/2} \preceq s_{n+1/2}$. La prima disuguaglianza segue facilmente dall'ipotesi induttiva. Per quanto riguarda la seconda, possiamo mostrare che $u_{n+1/2} \preceq s_{n-1/2}$ implica $u_{n+1/2} \preceq s_{n+1/2}$; poi ripetendo questo ragionamento $n - 1$ volte (che significa una induzione interna), otteniamo $u_{n+1/2} \preceq s_{1/2} \Rightarrow u_{n+1/2} \preceq s_{n+1/2}$, da cui la prova induttiva principale è completa, essendo banalmente vero che $u_{n+1/2} \preceq s_{1/2}$.

Ora assumiamo che $u_{n+1/2} \preceq s_{n-1/2}$. Per (2.18), $s_{n-1/2} \preceq t_n$ vale, dunque abbiamo

$$u_{n+1/2} \preceq t_n. \quad (2.21)$$

In più $u_n \preceq t_n$ per (2.19), allora (2.15), (2.16), (2.21) implicano $t_{n+1} \preceq u_{n+1/2}$. Richiamando (2.18), $s_n \preceq t_{n+1/2}$, e così $s_n \preceq u_{n+1/2}$. L'ultima disuguaglianza,

insieme con (2.15), (2.17) e (2.19), finalmente ci permettono di concludere che $u_{n+1/2} \preceq s_{n+1/2}$.

Per provare che $s_{n+1} \preceq u_{n+1}$, possiamo usare la stessa tecnica. In breve, si prova che $t_{1/2} \preceq u_{n+1/2}$ implica $t_{n-1/2} \preceq u_{n+1/2}$, e che $t_{n-1/2} \preceq u_{n+1/2}$ implica $t_{n+1/2} \preceq u_{n+1/2}$. per (2.19), $s_n \preceq u_n$, e per (2.15), (2.16) otteniamo $s_{n+1} \preceq u_{n+1}$.

Per concludere, essendo $\{s_n\}$ e $\{t_n\}$ convergenti entrambi a u^* , anche $\{u_n\}$ converge a u^* , per la normalità del cono $\mathcal{M}_+(\Omega, \mu)$. $\square$

L'equazione integrale (1.9)

$$u(x) = \int_0^{+\infty} \frac{\mathcal{K}(x, t)b(t)}{a(x)a(t)} \frac{1}{1+u(t)} d\mu(t). \quad (2.22)$$

vista nel capitolo precedente, è un caso particolare di (2.2) dove ψ è nulla,

$$g(u(t)) = \frac{1}{1+u(x)}, \quad k(x, t) = \frac{\mathcal{K}(x, t)b(t)}{a(x)a(t)}.$$

$\Omega = [0, +\infty)$. Questa particolare g è decrescente e $yg(y)$ strettamente crescente. Inoltre $\mathcal{K}(x, t)b(t)/(a(x)a(t))$ soddisfa l'ipotesi (1.12) con $\beta = b$. Si consideri ora lo splitting $A = E + F$ dove $A, E, F : \mathcal{M}_+(\Omega, \mu) \rightarrow \mathcal{M}_+(\Omega, \mu)$ sono gli operatori decrescenti definiti da

$$A(u)(x) = \int_0^\infty \frac{\mathcal{K}(x, t)b(t)}{a(x)a(t)} \frac{1}{1+u(t)} d\mu(t), \quad (2.23)$$

$$E(u)(x) = \int_0^x \frac{\mathcal{K}(x, t)b(t)}{a(x)a(t)} \frac{1}{1+u(t)} d\mu(t) \quad (2.24)$$

$$F(u)(x) = \int_x^\infty \frac{\mathcal{K}(x, t)b(t)}{a(x)a(t)} \frac{1}{1+u(t)} d\mu(t) \quad (2.25)$$

Allora, se l'equazione di Volterra $u = E(u) + h$ ammette un'unica soluzione in $\mathcal{M}_+((0, +\infty), \mu)$, valgono le ipotesi del Teorema 2.1.1 e del Teorema 2.1.2, e quindi la successione $\{u_n\}$ definita in (2.15) converge all'unica soluzione in $\mathcal{M}_+((0, +\infty), \mu)$.

La formulazione del problema nel caso infinito dimensionale ci permette di passare facilmente al caso finito dimensionale. Si prenda $\mathcal{M}(\Omega, \mu) = \mathbb{R}^m$ e $\mathcal{M}_+(\Omega, \mu) = \mathbb{R}_+^m$ (possiamo pensare ogni vettore $\mathbf{v}$ come una funzione che ad ogni i indice associa la i -esima componente: $i \rightarrow v_i$), e quale μ la misura

che conta i punti; essendo l'insieme vuoto l'unico insieme con misura nulla rispetto alla misura che conta i punti, tutti i “quasi ovunque” presenti nel caso infinito dimensionale scompaiono.

L'equazione di Hammerstein (2.22), si trasforma così nel sistema nonlineare

$$u_i = K_{ij} \frac{1}{1 + u_j} \quad i = 1, \dots, m \quad K_{ij} = \omega_j \frac{k(x_i, x_j)b(x_j)}{a(x_i)a(x_j)} \quad (2.26)$$

dove x_i e ω_i sono rispettivamente i nodi e i pesi di quadratura. Ora se scriviamo $\{K_{ij}\} = D + L + U$, dove D, L e U sono rispettivamente la diagonale, la parte triangolare inferiore e la parte triangolare superiore della matrice K , si ottengono 4 metodi rilassati dalle 4 scelte di E ed F :

$$\begin{aligned} (P_\omega) \quad E &= 0, \quad F = K; \quad (UP_\omega) \quad E = L, \quad F = U + D; \\ (J_\omega) \quad E &= D, \quad F = L + U; \quad (SOR) \quad E = L + D, \quad F = U; \end{aligned}$$

Questi metodi soddisfano le ipotesi del Teorema 2.1.1 e del Teorema 2.1.2, quindi convergono all'unica soluzione di (2.26).

2.2 Convergenza in L_+^∞ .

In questo paragrafo tratteremo la convergenza dello schema iterativo (2.1) a una soluzione dell'equazione (1.10) u^* nello spazio $L_+^\infty(\Omega, \mu)$ assumendo che valga l'ipotesi (1.13). Introduciamo ora alcuni risultati di carattere generale sugli spazi di Banach.

Teorema 2.2.1 *Sia $(X, P, \preceq)$ uno spazio di Banach ordinato da un cono normale, e $M : P \times P \rightarrow P$ un operatore mixed-monotone. Supponiamo che*

(H1) esistano $v \succ \theta$ e $c > 0$ tali che $\theta \preceq M(v, \theta) \preceq v$ e $M(\theta, v) \succeq cM(v, \theta)$;

(H2) per ogni $0 < a < b < 1$, esista $r = r(a, b) > 0$ tale che

$$M(\gamma s, t) \succeq \gamma(1 + r)M(s, \gamma t), \quad \text{per ogni } \gamma \in [a, b], \quad \theta \preceq t \preceq s \preceq v. \quad (2.27)$$

Allora M ha esattamente un punto fisso u^ in P , e $\theta \preceq u^* \preceq v$. Inoltre, costruendo ricorsivamente la successione*

$$s_{n+1} = M(s_n, t_n), \quad t_{n+1} = M(t_n, s_n), \quad n = 0, 1, 2, \dots, \quad (2.28)$$

abbiamo

$$\|s_n - u^*\| \rightarrow 0, \quad \|t_n - u^*\| \rightarrow 0, \quad \text{quando } n \rightarrow +\infty, \quad (2.29)$$

per ogni coppia iniziale $s_0, t_0 \in [\theta, v]$. In particolare, se $s_0 = \theta$, $t_0 = v$, sussiste la catena di disuguaglianze

$$\theta = s_0 \preceq s_1 \preceq \dots \preceq s_n \preceq \dots \preceq t_n \preceq \dots \preceq t_1 \preceq t_0 = A(\theta). \quad (2.30)$$

Teorema 2.2.2 Sia $(X, P, \preceq)$ uno spazio di Banach normalmente ordinato, e $A, E, F : P \rightarrow P$ operatori decrescenti, con $A = E + F$. Assumiamo che

(i) $A(\theta) \succ \theta$, ed esista $\epsilon_0 > 0$ tale che $A^2(\theta) \succeq \epsilon_0 A(\theta)$;

(ii) per ogni $0 < a < b < 1$ esista $\eta = \eta(a, b) > 0$ tale che

$$\gamma(1 + \eta)a(\gamma u) \preceq A(u), \quad \text{per ogni } \gamma \in [a, b], \quad u \in [\theta, A(\theta)]; \quad (2.31)$$

(iii) che l'equazione di punto fisso $u = E(u) + \beta$ abbia un'unica soluzione in P , per ogni $\beta \in [\theta, F(\theta)] \subseteq [\theta, A(\theta)]$.

Allora per ogni $\omega \in (0, 1]$ la successione $\{u_n\}$ definita da

$$\begin{cases} u_0 \in [\theta, A(\theta)] \\ u_{n+1/2} = E(u_{n+1/2}) + F(u_n), \\ u_{n+1} = (1 - \omega)u_n + \omega u_{n+1/2} \end{cases} \quad (2.32)$$

converge all'unico punto fisso, u^* , di A in P .

Dimostrazione. Fissato $\omega \in (0, 1]$, sia $M : P \times P \rightarrow P$ l'operatore

$$M(u, v) = (1 - \omega)u + \omega A(v), \quad \omega \in (0, 1]. \quad (2.33)$$

Posti

$$v = A(\theta), \quad c = \omega\epsilon_0, \quad r = \frac{\omega\epsilon_0\eta}{\omega\epsilon_0 + 1 - \omega}. \quad (2.34)$$

si ha che

$$\begin{aligned} M(v, \theta) &= (1 - \omega)A(\theta) + \omega A(\theta) = A(\theta), \\ M(\theta, v) &= \omega A^2(\theta) \succeq \omega\epsilon_0 A(\theta) = cM(v, \theta), \end{aligned} \quad (2.35)$$

e quindi l'ipotesi (H1) del Teorema 2.2.1 è verificata.

Riguardo (H2). per le posizioni (2.34), abbiamo

$$\omega(\eta - r)\epsilon_0 = (1 - \omega)r, \quad A(\gamma t) \succeq A^2(\theta) \succeq \epsilon_0 A(\theta) = \epsilon_0 \gamma s,$$

dunque, da (2.31) e (2.33)

$$\begin{aligned} M(\gamma s, t) &= (1 - \omega)\gamma s + \omega A(t) \succeq (1 - \omega)\gamma s + \omega\gamma(1 + \eta)A(\gamma t) \\ &\quad (1 - \omega)\gamma s + \omega\gamma(1 + r)A(\gamma t) + \omega\gamma(\eta - r)A(\gamma t) \\ &\succeq (1 - \omega)\gamma s + \omega\gamma(1 + r)A(\gamma t) + (1 - \omega)r\gamma s = \gamma(1 + r)M(s, \gamma t). \end{aligned} \quad (2.36)$$

Quindi per il Teorema 2.2.1 le successioni $\{s_n\}$, $\{t_n\}$ definite da

$$\begin{cases} t_0 = A(\theta) \\ t_{n+1/2} = E(u_n) + F(u_n) = A(u_n) \\ t_{n+1} = (1 - \omega)s_n + \omega t_{n+1/2}, \end{cases} \quad (2.37)$$

$$\begin{cases} s_0 = \theta \\ s_{n+1/2} = E(u_n) + F(u_n) = A(u_n) \\ s_{n+1} = (1 - \omega)t_n + \omega s_{n+1/2}, \end{cases} \quad (2.38)$$

convergono a u^* , l'unico punto fisso di M (e di P) in P . Con ragionamenti analoghi a quelli effettuati a quelli nel Teorema 2.1.2 concludiamo che pure $u_n \rightarrow u^*$. $\square$

Si consideri nuovamente l'equazione di Hammerstein (2.22) sostituendo k e g con le stesse funzioni. Assumiamo lo stesso splitting (2.25). Supponendo che valga l'ipotesi (1.13), e che l'equazione $u = Eu + h$ abbia una soluzione in $L_+^\infty([0, +\infty), \mu)$ per ogni $h \in [\theta, F(\theta)] \subseteq [\theta, A(\theta)]$, lo schema iterativo (2.32) del Teorema 2.2.2 converge all'unica soluzione di (2.22) in $L_+^\infty([0, +\infty), \mu)$.

Infatti se si considera $(L^\infty([0, +\infty), \mu), L_+^\infty([0, +\infty), \mu), \preceq)$ (spazio di Banach ordinato), gli operatori A , E e F , definiti da $L_+^\infty([0, +\infty), \mu)$ in sè stesso, sono chiaramente decrescenti. L'ipotesi (i) in relazione al Teorema 2.2.2 vale perchè $A(\theta) = \int_0^{+\infty} \mathcal{K}(x, t)b(t)/[a(t)a(x)] \cdot dt \prec 0$ per (1.13). Inoltre, se $u \in [\theta, A(\theta)]$, $\|u\|_\infty \leq \|A(\theta)\|_\infty = \sigma < \infty$ ancora per (1.13), e (ii) è verificata per η tale che

$$(1 + \eta)^{-1} = \left\{ \frac{\gamma(1 + u)}{1 + \gamma u}; (\gamma, u) \in [a, b] \times [0, \sigma] \right\}.$$

Ne deriva che l'uniforme convergenza dello schema iterativo (2.1) è garantita.

Concludiamo questo paragrafo con un cenno al caso finito dimensionale. Scegliendo opportunamente Ω e μ si ha che $L^\infty(\Omega, \mu) = \mathbb{R}^m$ e $L_+^\infty(\Omega, \mu) = \mathbb{R}_+^m$. Anche in questo caso, poiché nel caso finito l'ipotesi (1.13) è garantita, si ha la convergenza dei metodi $(UP)_\omega$, $(P)_\omega$, $(J)_\omega$ e $(SOR)_\omega$ alla soluzione del sistema (2.26) in $\mathbb{R}_+^m$.

Capitolo 3

Il metodo di Broyden

In questo capitolo mostreremo come risolvere dei sistemi non lineari utilizzando il metodo quasi-Newton di Broyden; in particolare ci interessano questioni di convergenza locale sotto le cosiddette ipotesi “standard”. Tratteremo inoltre alcune problematiche legate al test di arresto e all’implementazione di tali metodi numerici. Per una più estesa trattazione degli argomenti proposti in questo capitolo e per le dimostrazioni dei teoremi citati, rimandiamo il lettore a [11]. Il simbolo $\|\cdot\|$ indicherà una norma su $\mathbb{R}^N$, come anche la norma matriciale indotta.

Supponiamo dunque di dover approssimare una soluzione dell’equazione

$$F(u) = 0 \tag{3.1}$$

dove $F : \mathbb{R}^N \longrightarrow \mathbb{R}^N$. In seguito indicheremo con f_i la componente la i -esima componente di F . Se le componenti di f sono differenziabili in $u \in \mathbb{R}^N$ definiamo la *matrice jacobiana* $F'(u)$ con

$$F'(u)_{ij} = \frac{\partial f_i}{\partial u_j}(u).$$

Ipotesi 3.0.1 *Assumeremo su F le ipotesi standard.*

- (i) *L’equazione $F(u) = 0$ ha una soluzione u^* .*
- (ii) *$F' : \Omega \longrightarrow \mathbb{R}^{N \times N}$ è Lipschitziana (con costante di Lipschitz γ).*
- (iii) *$F'(u^*)$ è non singolare.*

Enunciamo ora un teorema fondamentale di analisi come segue.

Teorema 3.0.3 *Sia F differenziabile in un insieme aperto $\Omega \subset \mathbb{R}^N$ e sia $u^* \in \Omega$. Allora per ogni $u \in \Omega$ abbastanza vicini a u^**

$$F(u) - F(u^*) = \int_0^1 F'(u^* + t(u - u^*))(u - u^*) dt.$$

3.1 Tipi di convergenza

I metodi iterativi possono essere classificati per *velocità di convergenza*.

Definizione 3.1.1 *Sia $\{u_n\} \subset \mathbb{R}^N$ e $u^* \in \mathbb{R}^N$.*

(a) $u_n \rightarrow u^*$ q-quadraticamente se $u_n \rightarrow u^*$ ed esiste K tale che

$$\|u_{n+1} - u^*\| \leq K \|u_n - u^*\|^2.$$

(b) $u_n \rightarrow u^*$ q-superlinearmente con q-ordine $\alpha > 1$ se $u_n \rightarrow u^*$ ed esiste un K tale che

$$\|u_{n+1} - u^*\| \leq K \|u_n - u^*\|^\alpha$$

.

(c) $u_n \rightarrow u^*$ q-superlinearmente se

$$\lim_{n \rightarrow \infty} \frac{\|u_{n+1} - u^*\|}{\|u_n - u^*\|} = 0$$

(d) $u_n \rightarrow u^*$ q-linearmente con q-fattore $\sigma \in (0, 1)$ se

$$\|u_{n+1} - u^*\| \leq \sigma \|u_n - u^*\|$$

per n abbastanza grande.

Si noti che una successione convergente q-superlinearmente è anche convergente q-linearmente con q-fattore σ per ogni $\sigma > 0$. Una successione convergente q-quadraticamente è anche q-superlinearmente convergente con q-ordine 2. In generale un metodo convergente q-superlinearmente è preferibile ad che converge q-linearmente se il costo per ogni iterazione è lo stesso per entrambi i metodi. Spesso però un metodo che converge più lentamente, in termini di tipi di convergenza definiti sopra, può avere un costo per iterazione

così basso che risulta essere più efficiente.

In quanto a notazioni $B(r)$ è la palla di centro u^* e raggio r , x_n è l'ennesima iterata e e_n indica l'errore in quella iterata. s_n è la differenza tra due iterate successive, cioè

$$s_n = u_{n+1} - u_n.$$

Ci sarà utile il seguente lemma.

Lemma 3.1.1 *Assumiamo che l'ipotesi standard valga. Allora esiste $\delta > 0$ tale che per ogni $u \in B(u)$*

$$\|F'(u)\| \leq 2\|F'(u^*)\|, \quad (3.2)$$

$$\|F'(u)^{-1}\| \leq 2\|F'(u^*)^{-1}\| \quad (3.3)$$

$$\|F'(u^*)^{-1}\|^{-1}\|e\|/2 \leq \|F(u)\| \leq 2\|F'(u^*)\|\|e\| \quad (3.4)$$

Per la dimostrazione si veda [11].

3.2 Metodo di Newton

In questo paragrafo descriveremo brevemente il metodo di Newton per equazioni non lineari considerando come già anticipato, questioni di convergenza locale. Se u_c è l'iterata corrente, il metodo di Newton produce quale nuova approssimazione della soluzione

$$u_+ = u_c - F'(u_c)^{-1}F(u_c) \quad (3.5)$$

Una condizione sufficiente per la convergenza locale del metodo di Newton segue dal Lemma 3.1.1.

Teorema 3.2.1 *Supponiamo che valgano le ipotesi standard. Allora esistono $K > 0$ e $\delta > 0$ tali che se $u_c \in B(\delta)$ l'iterata di Newton data da (3.5) a partire da u_c soddisfa*

$$\|e_+\| \leq K\|e_c\|^2 \quad (3.6)$$

Dimostrazione. Sia δ abbastanza piccolo tale che valgano le conclusioni del Lemma 3.1.1. Per il Teorema 3.0.3

$$e_+ = e_c - F'(u_c)^{-1}F(u_c) = F'(u_c)^{-1} \int_0^1 (F'(c_c) - f'(u^* + te_c))e_c dt.$$

per il Lemma 3.1.1 e la Lipschitzianità di F'

$$\|e_+\| \leq (2\|F'(u^*)^{-1}\|)\gamma\|e_c\|^2/2$$

Questo completa la dimostrazione del teorema con $K = \gamma\|F'(u^*)^{-1}\|$. $\square$

Teorema 3.2.2 *Supponiamo valgano le ipotesi standard. Allora esiste un δ tale che se $u_0 \in B(\delta)$ allora l'iterazione di Newton*

$$u_{n+1} = u_n - F'(u_n)^{-1}F(u_n)$$

converge q-quadraticamente a u^ .*

Dimostrazione. Sia δ abbastanza piccolo tale che valga la conclusione del 3.2.1. Si prenda δ abbastanza piccolo tale che $K\delta = \eta < 1$. Allora se $n \geq 0$ e $u_n \in B(\delta)$, il Teorema 3.2.1 implica che

$$\|e_{n+1}\| \leq K\|e_n\|^2 \leq \eta\|e_n\| < \|e_n\| \quad (3.7)$$

e dunque $u_{n+1} \in B(\eta\delta) \subset B(\delta)$. Dunque se $u_n \in B(\delta)$ possiamo continuare l'iterazione. Dato che $u_0 \in B(\delta)$ per ipotesi, l'intera successione $\{u_n\} \subset B(\delta)$. (3.7) allora implica che $u_n \rightarrow u^*$ q-quadraticamente. $\square$

3.3 Test di uscita

Abbiamo deciso di usare come test di uscita lo step relativo, cioè di terminare le iterazioni quando

$$\frac{\|u_{n+1} - u_n\|}{\|u_{n+1}\|} \leq \text{tol}$$

Questo criterio è basato sul teorema 3.2.1. Infatti

$$\begin{aligned} \|u_n - u^*\| &= \|u_n - u_{n+1} + u_{n+1} - u^*\| \leq \|u_n - u_{n+1}\| + \|u_{n+1} - u^*\| \\ &\leq \|u_n - u_{n+1}\| + K\|u_n - u^*\|^2 \end{aligned}$$

Dunque abbiamo

$$\|u_n - u^*\| \leq \|u_{n+1} - u_n\| + K\|u_n - u^*\|^2$$

Allora, vicino alla soluzione, lo $\|u_{n+1} - u_n\|$ e $\|u_n - u^*\|$ sono essenzialmente gli stessi. useremo quindi come test di uscita per le iterazioni nel metodo di Newton il seguente test:

$$\frac{\|u_n - u_{n-1}\|}{\|u_n\|} \leq \text{tol}$$

dove tol è la tolleranza sotto la quale vogliamo tenere l'errore relativo.

3.3.1 Implementazione del metodo di Newton

Consideriamo la generica iterazione (3.5) del metodo di Newton. Per ottenere u_+ si procede calcolando la “correzione s_c quale soluzione del sistema lineare

$$F'(u_c)s_c = -F(u_c),$$

dopo di che $x_+ = x_c + s_c$. Negli esempi considerati la matrice Jacobiana è densa e la sua valutazione e fattorizzazione risultano onerose. Infatti fattorizzare F' nel caso denso costa $O(N^3)$ flops e la valutazione di F' per differenze finite è da aspettarsi che costi N volte il costo della valutazione di F , perchè ogni colonna in F' ha bisogno di una valutazione di F per formare l'approssimazione della differenza. Nella descrizione dell'algoritmo di Newton noi usiamo la fattorizzazione LU di F' , anche se altre fattorizzazioni possono essere considerate. Fatte queste considerazioni, fissate l'iterata iniziale, la funzione F e la tolleranza per il test di uscita tol , il metodo di Newton è descritto dal seguente pseudo-codice.

Algoritmo 3.3.1 newton (u, F, tol)

1. $r_0 = \frac{\|u_n - u_{n-1}\|}{\|u_n\|}$
2. Do while $r_0 < tol$
 - (a) calcola $F'(u_n)$
 - (b) Fattorizza $F'(u_n) = LU$
 - (c) Risolvi $LU s_n = -F(u)$
 - (d) $u_{n+1} = u_n + s_n$

Come abbiamo visto, portare a termine i punti 2a e 2b è la parte più costosa dell'algoritmo. Introduciamo ora un metodo *quasi-Newton*, il *metodo di Broyden*, che pur mantenendo la struttura del metodo di Newton, alleggerisce il calcolo dei punti 2a e 2b.

3.4 Il metodo di Broyden.

Abbiamo nella sezione precedente che uno dei maggiori problemi del metodo di Newton consiste nel calcolare la matrice Jacobiana e risolvere un certo sistema lineare. Nel metodo di Broyden la matrice Jacobiana viene approssimata a basso costo, mantenendo sotto opportune ipotesi (condizione di Dennis-Moré) la convergenza superlineare. Se x_c e B_c sono le approssimazioni correnti della soluzione e dello Jacobiano, allora

$$u_+ = u_c - B_c^{-1}F(u_c). \quad (3.8)$$

Nel metodo di Broyden quale approssimazione di $JF(x_n)$ si considera

$$B_+ = B_c + \frac{(y - B_c s)s^T}{s^T s} = B_c + \frac{F(u_+)s^T}{s^T s}, \quad (3.9)$$

ove $y = F(u_+) - F(u_c)$ e $s = u_+ - u_c$.

Per verificare la convergenza superlineare dei metodi quasi-Newton si usa generalmente la *condizione di Dennis-Moré*

$$\lim_{n \rightarrow \infty} \frac{\|E_n s_n\|}{\|s_n\|} = 0, \quad (3.10)$$

dove $s_n = u_{n+1} - u_n$ e E_n è l'errore sull'approssimazione della matrice Jacobiana.

Per quanto concerne la convergenza locale risulta fondamentale il seguente

Teorema 3.4.1 *Supponiamo valgano le ipotesi standard; siano $\{B_n\}$ una successione di matrici $N \times N$ non singolari, $u_0 \in \mathbb{R}^N$ dato e $\{u_n\}_{n=1}^\infty$ data da (3.8). Assumiamo che $u_n \neq u_*$ per ogni n . Allora $u_n \rightarrow u^*$ q-superlinearmente se e solo se $u_n \rightarrow u^*$ soddisfacendo al condizione (3.10).*

3.4.1 Convergenza del metodo di Broyden.

Iniziamo subito con una definizione.

Definizione 3.4.1 *Diciamo che la successione di Broyden $(\{u_n\}, \{B_n\})$ per i dati (F, u_0, B_0) esiste se F è definita su u_n per ogni n , B_n è non singolare per ogni n , e u_{n+1} e B_{n+1} possono essere calcolati da u_n e B_n usando (3.8) e (3.9).*

Mostreremo che se valgono le assunzioni standard e se u_0 e B_0 sono sufficientemente buone approssimazioni di u^* e $F'(u^*)$, allora la successione di Broyden esiste per i dati (F, u_0, B_0) , e che le iterate di Broyden convergono q-superlinearmente a u^* . Lo sviluppo della prova è complicato. Prima proveremo un risultato di *bounded deterioration*; mostreremo poi che se i dati iniziali sono abbastanza buoni, la crescita degli errori può essere limitata e dunque esiste la successione di Broyden e le iterate di Broyden convergono a u^* almeno q-quadraticamente. Infine verifichiamo la condizione (3.10) di Dennis-Moré, che, insieme con la convergenza q-lineare provata precedentemente, completa la prova della convergenza locale q-superlineare.

Bounded Deterioration.

Teorema 3.4.2 *Supponiamo valgano le ipotesi standard. Siano $u_c \in \Omega$ e una matrice B_c dati. Assumiamo che*

$$u_+ = u_c - B_c^{-1}F(u_c) = u_c + s \in \Omega$$

e B_+ sia dato da (3.9). Allora

$$\|E_+\|_2 \leq \|E_c\|_2 + \gamma(\|e_c\|_2 + \|e_+\|_2)/2$$

Convergenza q-lineare locale

Teorema 3.4.3 *Supponiamo valgano le ipotesi standard. Sia $r \in (0, 1)$ dato. Allora esistono δ e δ_B tali che se $u_0 \in B(\delta)$ e $\|E_0\|_2 < \delta_B$, la successione di Broyden per i dati (F, u_0, B_0) esiste e $u_n \rightarrow u^*$ q-linearmente con q-fattore al massimo r .*

Verifica della condizione di Dennis-moré.

Teorema 3.4.4 *Supponiamo che valgano le ipotesi standard. Allora esistono δ e δ_B tali che se $u_0 \in B(\delta)$ e $\|E_0\|_2 < \delta_B$, la successione di Broyden per i dati (F, u_0, B_0) esiste e $u_n \rightarrow u^*$ q-superlinearmente.*

Per la dimostrazione di questo teorema basta prendere δ e δ_B tali che valgano le conclusioni del Teorema 3.4.4 e, per il Teorema 3.4.1, far valere la condizione (3.10) di Dennis-Moré.

3.4.2 Implementazione del metodo di Broyden

Iniziamo mostrando come l'approssimazione iniziale di $F'(u^*)$, B_0 può essere considerato come un preconditionatore e incorporato nella definizione di F .

Lemma 3.4.1 *Assumiamo che $(\{u_n\}, \{B_n\})$ sia la successione di Broyden per i dati (F, u_0, B_0) . Allora la successione di Broyden $(\{\xi_n\}, \{C_n\})$ esiste per i dati $(B_0^{-1}F, u_0, I)$ e*

$$u_n = y_n, \quad B_n = B_0 C_n \quad (3.11)$$

per tutti gli n .

Dimostrazione. Procediamo per induzione. Equazione (3.11) vale per definizione quando $n = 0$. Se (3.11) vale per un dato n , allora

$$\begin{aligned} y_{n+1} &= y_n - C_n^{-1} C_0^{-1} F(y_n) \\ &= u_n - (C_0 C_n)^{-1} F(u_n) = u_n - B_n^{-1} F(u_n) \\ &= u_{n+1}. \end{aligned} \quad (3.12)$$

Ora $s_n = u_{n+1} - u_n = y_{n+1} - y_n$ per (3.12). Allora

$$\begin{aligned} B_{n+1} &= B_n + \frac{F(u_{n+1})}{\|s_n\|_2^2} \\ &= B_0 C_n + B_0 \frac{B_0^{-1} F(y_{n+1})}{\|s_n\|_2^2} = B_0 C_{n+1}. \end{aligned}$$

Questo completa la dimostrazione. $\square$

La conseguenza di questo lemma per l'implementazione è che possiamo assumere $B_0 = I$. L'idea è di memorizzare B_n^{-1} in un modo compatto e facile da usare. La base per questo è la seguente formula.

Proposizione 3.4.1 *Sia B una matrice $N \times N$ non singolare e siano $u, v \in \mathbb{R}^N$. Allora $B + uv^T$ è invertibile se e solo se $1 + v^T B^{-1} u \neq 0$. In questo caso*

$$(B + uv^T)^{-1} = \left(I - \frac{(B^{-1}u)v^T}{1 + v^T B^{-1}u} \right) B^{-1} \quad (3.13)$$

L'espressione (3.13) è chiamata la *formula di Sherman-Morrison*. Per quanto riguarda l'aggiornamento di una successione di Broyden $\{B_n\}$, abbiamo per $n \geq 0$

$$B_{n+1} = B_n + u_n v_n^T,$$

dove

$$u_n = F(u_{n+1})/\|s_n\| \text{ e } v_n = s_n/\|s_n\|.$$

Ponendo

$$w_n = (B_n^{-1})/(1 + v_n^T B_n^{-1} u_n)$$

vediamo che, se $B_0 = I$,

$$\begin{aligned} B_n^{-1} &= (I - w_{n-1} v_{n-1}^T)(I - w_{n-2} v_{n-2}^T) \dots (I - w_0 v_0^T) \\ &= \prod_{j=0}^{n-1} (I - w_j v_j^T) \end{aligned} \quad (3.14)$$

Per il fatto che il prodotto vuoto di matrici è l'identità, (3.14) è valido per $n \geq 0$.

Allora l'azione di B_n^{-1} su $F(u_n)$ (cioè, il calcolo del passo di Broyden) può essere calcolato dai $2n$ vettori $\{w_j, v_j\}_{j=0}^{n-1}$ al costo di $O(Nn)$ operazioni a virgola mobile. Inoltre il passo di Broyden per la seguente iterazione è

$$s_n = -B_n^{-1} F(u_n) = - \sum_{j=0}^{n-1} (I - w_j v_j^T) F(u_n). \quad (3.15)$$

Siccome il prodotto

$$\prod_{j=0}^{n-2} (I - w_j v_j^T) F(u_n)$$

deve essere calcolato come parte del calcolo di w_{n-1} , possiamo combinare il calcolo di w_{n-1} e s_n come segue

$$\begin{aligned} w &= \prod_{j=0}^{n-2} (I - w_j v_j^T) F(u_n) \\ w_{n-1} &= C_w w \text{ dove } C_w = (\|s_{n-1}\|_2 + v_{n-1}^T w)^{-1} \\ s_n &= -(I - w_{n-1} v_{n-1}^T) w \end{aligned} \quad (3.16)$$

Si può proseguire ed eliminare il bisogno di memorizzare la successione $\{w_n\}$. (3.16) implica che per $n \geq 1$

$$\begin{aligned} s_n &= -(w - C_w w(v_{n-1}^T w)) = -w(1 - C_w(C_w^{-1} - \|s_{n-1}\|_2)) \\ &= -C_w w \|s_{n-1}\|_2 = -\|s_{n-1}\|_2 w_{n-1}. \end{aligned} \quad (3.17)$$

Dunque, per $n \geq 0$,

$$w_n = -s_{n+1}/\|s_n\|_2. \quad (3.18)$$

Abbiamo perciò bisogno di memorizzare la successione degli steps $\{s_n\}$ e le loro norme per costruire la successione $\{w_n\}$. Infatti possiamo scrivere (3.14) come

$$B_n^{-1} = \prod_{j=0}^{n-1} \left(I + \frac{s_{j+1} s_j^T}{\|s_j\|_2^2} \right). \quad (3.19)$$

Non possiamo usare direttamente (3.19) e (3.15) per calcolare s_{n+1} perché s_{n+1} compare da entrambe le parti dell'equazione

$$\begin{aligned} s_{n+1} &= - \left(I - \frac{s_{n+1} s_n^T}{\|s_n\|_2^2} \right) \prod_{j=0}^{n-1} \left(I - \frac{s_{j+1} s_j^T}{\|s_j\|_2^2} \right) F(u_{n+1}) \\ &= - \left(I - \frac{s_{n+1} s_n^T}{\|s_n\|_2^2} \right) B_n^{-1} F(u_{n+1}). \end{aligned} \quad (3.20)$$

Risolviamo invece (3.20) ed otteniamo

$$s_{n+1} = \frac{B_n^{-1} F(u_{n+1})}{1 + s_n^T B_n^{-1} F(u_{n+1}) / \|s_n\|_2^2}. \quad (3.21)$$

Per la Proposizione 3.4.1, il denominatore in (3.21)

$$1 + s_n^T B_n^{-1} F(u_{n+1}) / \|s_n\|_2^2 = 1 + v_n^T B_n^{-1} u_n$$

è diverso da zero quando B_{n+1} è non singolare.

Capitolo 4

Affrontiamo ora la risoluzione numerica dell'equazione quadratica

$$a(x)\phi(x) + \phi(x) \int_{\Omega} \mathcal{K}(x, t)\phi(t) d\mu(t) = b(x), \quad x \in \Omega. \quad (4.1)$$

Sappiamo che attraverso la trasformazione

$$\phi = \frac{b}{a} \frac{1}{1 + u}, \quad (4.2)$$

diventa l'equazione di Hammerstein

$$u(x) = \int_{\Omega} \frac{\mathcal{K}(x, t)b(t)}{a(x)a(t)} \frac{1}{1 + u(t)} d\mu(t), \quad x \in \Omega. \quad (4.3)$$

Calcoleremo una approssimazione $\tilde{\phi}$ della soluzione ϕ attraverso la risoluzione numerica di (4.3), cioè calcoleremo una soluzione approssimata $\tilde{u}$ di (4.3) e per la posizione (4.2) arriveremo a $\tilde{\phi}$. (4.2) ci garantisce che la soluzione dell'equazione quadratica che troviamo stia in $L^1_+(\Omega, \mu)$. Considereremo (4.3) due differenti formulazioni dell'equazione (4.3), una con nucleo di tipo Boltzmann e un'altra con nucleo di tipo Chandrasekhar. Precisamente avremo:

(i) nel caso di Boltzmann

$$\frac{\mathcal{K}(x, t)b(t)}{a(x)a(t)} = \lambda \frac{t^{1-p}e^{-(t)^2}((x+t)^{q+2} - |x-t|^{q+2})}{2(q+2)x^{p+2}},$$

dove $a(x) = x^p$, $b(x) = \lambda x^2 e^{-x^2}$ e l'insieme di integrazione $\Omega = (0, +\infty)$;

- (ii) nel caso di Chandrasekhar sarà $\mathcal{K}(x, t) = x/(x + t)$ con $\mathcal{K}(0, 0) = 1$. In questo caso $a(x) = 1$, $b(x) = 1$ per ogni x nell'intervallo di integrazione $(0, 1)$.

I solutori implementati adottano una tecnica “multilivello”. Fissato un livello, cioè fissato un certo numero di nodi e di pesi per una opportuna formula di quadratura, l'equazione (4.3) viene discretizzata e si trasforma nel sistema non lineare

$$(u^l)_i = \sum_{j=0}^{N^l} w_j k(x_i, x_j) \frac{1}{1 + (u^l)_j}, \quad i = 0, \dots, N^l, \quad (4.4)$$

con $K_{ij} = w_j k(x_i, x_j)$ (k è la parte lineare dell'equazione di Hammerstein) e dove $\{x_j\}$, $\{w_j\}$ sono rispettivamente i nodi e i pesi di quadratura. A questo punto, con un metodo iterativo, verrà approssimata la soluzione di (4.4) con u_m^l , che è l'approssimazione dopo m iterazioni della soluzione di (4.4) al livello l . Nel caso in cui ϕ_m^l , calcolato a partire da u_m^l attraverso (4.2) non approssimi sufficientemente bene la soluzione di (4.1), la discretizzazione viene affinata al livello $l + 1$ ($N^{l+1} > N^l$, ad esempio $N^{l+1} = 2N^l$) dando origine ad un nuovo sistema non lineare. Si procede in questo modo finché l'approssimazione trovata è abbastanza buona, cioè il vettore ϕ_m^l calcolato a partire da u_m^l con (4.2) approssima bene la soluzione dell'equazione quadratica (4.1).

Abbiamo implementato tre solutori, due che si basano su metodi di tipo Picard, (Updated Picard e Picard semplice), e uno che si basa sul metodo di quasi-Newton di Broyden. Sono tre solutori multilivello nel senso prima precisato, che a partire da un vettore iniziale, generalmente il vettore nullo, approssimano la soluzione dell'equazione, passando attraverso vari livelli. Per permettere il passaggio dell'approssimazione da un livello ad un altro, il che vuol dire da un insieme di pesi e nodi di quadratura ad un altro, abbiamo bisogno di una formula interpolatoria. Nei solutori qui implementati ci siamo serviti della formula di interpolazione di Nyström [1]. Dato un vettore $\{(u^l)_i\}_{0 \leq i \leq N^l}$ che risolve (4.4) possiamo definire per ogni x che stà nell'intervallo di integrazione

$$u^l(x) = \sum_{j=0}^{N^l} w_j^l k(x, x_j^l) \frac{1}{1 + u_j^l} \quad j = 0, \dots, N^l, \quad (4.5)$$

dove $\{x_i^l\}_{0 \leq i \leq N^l}$ e $\{w_i^l\}_{0 \leq i \leq N^l}$ sono i nodi e pesi di quadratura del livello l . Si tratta sicuramente di una formula interpolatoria in quanto

$$u^l(x_i) = \sum_{j=0}^{N^l} w_j^l k(x_i, x_j^l) \frac{1}{1 + u_j^l} = u_i^l, \quad j = 0, \dots, N^l,$$

essendo u_i^l soluzione di (4.4). Dunque se abbiamo una approssimazione $\{(\tilde{u}^l)_i\}_{0 \leq i \leq N^l}$ della soluzione di (4.4) al livello l ed è richiesto di continuare le iterazioni al livello $l + 1$, interpoliamo $\tilde{u}^l$ con la formula di Nyström

$$u_0^{l+1}(x_i^{l+1}) = \sum_{j=0}^{N^l} w_j^l k(x_i^{l+1}, x_j^l) \frac{1}{1 + (\tilde{u}^l)_j} \quad j = 0, \dots, N^{l+1}, \quad (4.6)$$

e questo sarà il vettore da cui si partirà per calcolare una approssimazione della soluzione di (4.4) al livello $l + 1$.

Nelle stime degli errori che discuteremo in seguito, applicheremo la formula di Nyström anche a ϕ^l

$$\tilde{\phi}^l(x) = \sum_{j=0}^{N^l} w_j^l k(x, x_j^l) \frac{1}{1 + \tilde{u}_j^l}. \quad (4.7)$$

Per quanto riguarda i nodi e i pesi di quadratura, abbiamo scelto la formula che genera i nodi e i pesi di Fejer [7]. È una formula interpolatoria aperta, fatta sui nodi di Chebyshev nell'intervallo $(-1, 1)$ (perché su questi nodi non compare il fenomeno di Runge) ed è convergente su tutte le funzioni continue, anche quelle con singolarità integrabili agli estremi. Le formule che generano i nodi e i pesi per l'intervallo $(-1, 1)$ sono

$$\hat{x}_i^l = \cos((2(i + 1) - 1)\pi/(2(N^l + 1))), \quad i = 0, \dots, N^l - 1; \quad (4.8)$$

$$\hat{w}_i^l = \frac{2}{N + 1} \left(1 - 2 \sum_{s=1}^{[(N+1)/2]} \frac{\cos(2s\theta_i)}{4s^2 - 1} \right), \quad i = 0, \dots, N^l - 1, \quad (4.9)$$

dove $\theta_i = \arccos(\hat{x}_i)$.

Poiché generalmente l'intervallo di integrazione Ω è diverso da $(-1, 1)$, abbiamo bisogno di una funzione che distribuisca i nodi e pesi all'interno di Ω . Se l'intervallo di integrazione $\Omega = (a, b)$ è finito, i nodi diventano $x_i^l = [(b - a)/2](\hat{x}_i^l - 1) + b$ mentre i pesi $w_i^l = \tilde{w}_i^l(b - a)/2$. L'altra possibile

alternativa che prenderemo in considerazione è il caso in cui l'intervallo di integrazione sia $\Omega = (a, +\infty)$, con $a \geq 0$. I nodi e i pesi diventeranno rispettivamente $x_i^l = \Upsilon(\hat{x}_i^l)$ e $w_i^l = \hat{w}_i^l \Upsilon'(\hat{x}_i^l)$, $i = 0, \dots, N^l$, dove $\Upsilon(x) = 2(a + 1 + x)/(1 - x)$. Si noti che $\lim_{x \rightarrow 1} \Upsilon(x) = +\infty$ e $\Upsilon(-1) = a$.

Passiamo ora alla descrizione della struttura iterativa dei solutori. Come abbiamo detto, risolviamo l'equazione di Hammerstein (4.3) per ottenere una soluzione dell'equazione quadratica (4.1) attraverso la trasformazione (4.2). Dunque per decidere quando terminare le iterazioni valuteremo $\|\phi_m^l - \{\phi(x_i^l)\}\|_1$. La norma $\|\cdot\|_1$, per come è definita, approssima la norma in L^1 . Infatti $\|\{\phi(x_i^l)\}\|_1 = \sum_{i=0}^{N^l} w_i^l \|\phi(x_i^l)\| \approx \|\phi\|_{L^1}$. L'errore assoluto sarà dunque $\|\phi_m^l - \{\phi(x_i^l)\}\|_1 = \sum_{i=0}^{N^l} w_i^l |(\phi_m^l)_i - \phi(x_i^l)|$. Servendosi della trasformazione

$$(\phi^l)_i = \frac{b(x_i^l)}{a(x_i^l)} \frac{1}{1 + (u^l)_i} \quad (4.10)$$

otteniamo

$$\|\phi_m^l - \{\phi(x_i^l)\}\|_1 = \sum_{i=0}^{N^l} w_i^l \frac{b(x_i^l)}{a(x_i^l)} \frac{|\{u(x_i^l)\} - (u^l)_i|}{(1 + (u^l)_i)(1 + \{u(x_i^l)\})}, \quad (4.11)$$

e

$$\|\phi_m^l\|_1 = \sum_{i=0}^{N^l} w_i^l \frac{b(x_i^l)}{a(x_i^l)} \left| \frac{1}{1 + (u_m^l)_i} \right|; \quad (4.12)$$

ci serviremo di (4.12) e (4.11) nell'implementazione dei solutori. Supponiamo ora che, ad un livello l , il solutore stia approssimando con ϕ_m^l la soluzione ϕ^l attraverso il calcolo di u_m^l , approssimazione di u^l , che è la soluzione di (4.4). Decidiamo di arrestare il processo iterativo del livello l ad un certo m tale che

$$\frac{\|\phi_m^l - \{\phi(x_i^l)\}\|_1}{\|\{\phi(x_i^l)\}\|_1} \leq \epsilon^l, \quad (4.13)$$

dove ϵ^l è una approssimazione di $\|\tilde{\phi}^l - \phi(x_i^l)\|_1 / \|\{\phi(x_i^l)\}\|_1$, che è l'errore dovuto alla discretizzazione di (4.3) al livello l , alla risoluzione del sistema (4.4) e alla formula interpolatoria di Nyström. Ad ogni livello chiederemo una precisione pari all'errore $\|\tilde{\phi}^l - \phi(x_i^l)\|_1 / \|\{\phi(x_i^l)\}\|_1$ e questo è il meglio che si possa fare a quel determinato livello, in quanto, anche se si chiedesse una precisione maggiore, dominerebbe l'errore $\|\tilde{\phi}^l - \phi(x_i^l)\|_1 / \|\{\phi(x_i^l)\}\|_1$. Il problema stà ora nel calcolo di ϵ^l , tenendo conto che siamo interessati all'ordine

di grandezza dell'errore di discretizzazione. Facciamo vedere che

$$\|\phi_2^{l+1} - \phi_0^{l+1}\|_1 \approx \|\phi^l - \{\phi(x_i^l)\}\|_1, \quad (4.14)$$

e che dunque è un'approssimazione ragionevole dell'errore assoluto al livello l .

Abbiamo deciso di considerare $\|\phi_2^{l+1} - \phi_0^{l+1}\|_1$ e non $\|\phi_1^{l+1} - \phi_0^{l+1}\|_1$ perchè in questo modo la stima dell'errore di discretizzazione risulta essere più affidabile. Questo è giustificato dal fatto che per esempio nel metodo di Broyden l'iterata u_1^l viene calcolata approssimando l'inversa dello jacobiano con la matrice identità e ciò non garantisce una sufficiente precisione delle stime fatte. Quanto detto è stato notato anche sperimentalmente testando i programmi e ci ha spinto ad addottare questa stima in tutti i nostri programmi.

Premettiamo una semplice osservazione che ci sarà utile.

Osservazione. Dati due vettori $\mathbf{v}$, $\mathbf{w}$ tali che $\|\mathbf{w}\| = \theta\|\mathbf{v}\|$, $0 < \theta < 1$, otteniamo che

$$\|\mathbf{v}\|(1 - \theta) \leq \|\mathbf{v} + \mathbf{w}\| \leq \|\mathbf{v}\|(1 + \theta).$$

Si può notare che se θ è “piccolo”, quel che domina è l'ordine di grandezza di $\|\mathbf{v}\|$. Questo vale già a partire da $\theta = 1/3$. In questi casi possiamo scrivere $\mathbf{v} + \mathbf{w} \approx \mathbf{v}$.

Tornando al nostro problema,

$$\phi_2^{l+1} - \phi_0^{l+1} = \phi_2^{l+1} - \phi^{l+1} + \phi^{l+1} - \phi_0^{l+1} \approx \phi^{l+1} - \phi_0^{l+1}$$

essendo ϕ_2^{l+1} più vicino a ϕ^{l+1} di ϕ_0^{l+1} ,

$$= \phi^{l+1} - \{\phi(x_i^{l+1})\} + \{\phi(x_i^{l+1})\} - \phi_0^{l+1} \approx \{\phi(x_i^{l+1})\} - \phi_0^{l+1}$$

essendo ancora una volta ϕ^{l+1} molto più vicino a $\{\phi(x_i^{l+1})\}$ di ϕ_0^{l+1} . Otteniamo dunque

$$\phi_2^{l+1} - \phi_0^{l+1} \approx \{\phi(x_i^{l+1})\} - \phi_0^{l+1} = \{\phi(x_i^{l+1})\} - \{\tilde{\phi}^l(x_i^{l+1})\},$$

dove l'ultima uguaglianza è giustificata dalla formula interpolatoria di Nyström applicata a ϕ^l (4.7). Ora

$$\|\{\tilde{\phi}^l(x_i^{l+1})\} - \{\phi(x_i^{l+1})\}\|_1 = \sum_{i=0}^{N^{l+1}} w_i^{l+1} |\tilde{\phi}^l(x_i^{l+1}) - \phi(x_i^{l+1})| \approx \int_{\Omega} |\tilde{\phi}^l(t) - \phi(t)| dt,$$

dove l'ultimo termine ingloba contemporaneamente sia l'errore della risoluzione del sistema (4.4), sia l'errore della formula interpolatoria di Nyström applicata per ottenere la funzione $\tilde{\phi}^l(t)$. Ma anche

$$\sum_{i=0}^{N^l} w_i^l |\tilde{\phi}^l(x_i^l) - \phi(x_i^l)| \approx \int_{\Omega} |\tilde{\phi}^l(t) - \phi(t)| dt,$$

essendo il primo termine ancora una approssimazione dell'integrale, anche se fatta su un numero minore di nodi e pesi di quadratura. Sapendo che

$$\sum_{i=0}^{N^l} w_i^l |\tilde{\phi}^l(x_i^l) - \phi(x_i^l)| = \|\tilde{\phi}^l - \{\phi(x_i^l)\}\|_1$$

otteniamo che

$$\|\phi_2^{l+1} - \phi_0^{l+1}\|_1 \approx \|\tilde{\phi}^l - \{\phi(x_i^l)\}\|_1.$$

Possiamo scrivere

$$\tilde{\phi}^l(x_i^l) - \phi(x_i^l) = (\tilde{\phi}^l(x_i^l) - \phi^l(x_i^l)) + (\phi^l(x_i^l) - \{\phi(x_i^l)\}); \quad (4.15)$$

siccome per (4.13) il primo addendo del secondo membro di (4.15) è al massimo dello stesso ordine di grandezza del secondo addendo, che è in norma $\|\cdot\|_1$ l'errore di discretizzazione al livello l insieme con l'errore della formula di Nyström, possiamo concludere che

$$\|\phi_2^{l+1} - \phi_0^{l+1}\|_1 \approx \|\{\tilde{\phi}^l(x_i^l)\} - \{\phi(x_i^l)\}\|_1 \approx \|\{\phi^l(x_i^l)\} - \{\phi(x_i^l)\}\|_1, \quad (4.16)$$

che è proprio quello che volevamo giustificare.

Assumeremo dunque assumeremo

$$\epsilon^l = \frac{\|\phi_2^{l+1} - \phi_0^{l+1}\|_1}{\|\{\phi(x_i^{l+1})\}\|_1}. \quad (4.17)$$

Come si vede da (4.17) riusciamo a calcolare ϵ^l solo al livello $l+1$. Avremo quindi bisogno di stimare ϵ^l già al livello l . Per fare questo moltiplichiamo ϵ^{l-1} per il rapporto $\epsilon^{l-1}/\epsilon^{l-2}$, che approssima di quanto è diminuito l'errore di discretizzazione al passaggio dal livello $l-2$ al livello $l-1$. In questo modo $(\epsilon^{l-1})^2/\epsilon^{l-2}$ stima ϵ^l . Per ottenere ciò al livello l , tenendo conto di (4.17), calcoleremo

$$\epsilon^l \approx \left(\frac{\|\phi_2^l - \phi_0^l\|_1}{\|\{\phi(x_i^l)\}\|_1} \bigg/ \frac{\|\phi_2^{l-1} - \phi_0^{l-1}\|_1}{\|\{\phi(x_i^{l-1})\}\|_1} \right) \frac{\|\phi_2^l - \phi_0^l\|_1}{\|\{\phi(x_i^l)\}\|_1} \approx$$

$$\frac{\|\phi_2^l - \phi_0^l\|_1^2}{\phi_2^{l-1} - \phi_0^{l-1}} \frac{1}{\|\{\phi(x_i^l)\}\|_1} = \eta^l. \quad (4.18)$$

Osserviamo che nell'implementazione del test (4.13) stimeremo ϵ^l con il massimo tra η^l e tol data, perchè non sarebbe ragionevole chiedere ad un determinato livello una precisione maggiore di quella richiesta dal problema. (4.13) diventa

$$\frac{\|\phi_m^l - \{\phi(x_i^l)\}\|_1}{\|\{\phi(x_i^l)\}\|_1} \leq \max\{\eta^l, tol\}. \quad (4.19)$$

Inoltre affinché abbiano senso le stime fatte il solutore potrà utilizzare il test (4.19) non prima del livello $l = 2$, supponendo che il livello iniziale sia $l = 0$. Per i livelli 0 e 1 il test sarà semplicemente

$$\frac{\|\phi_m^l - \{\phi(x_i^l)\}\|_1}{\|\{\phi(x_i^l)\}\|_1} \leq tol. \quad (4.20)$$

Questa scelta risulta essere legittima poiché ai primi livelli usualmente l'equazione viene discretizzata su pochi punti e dunque il costo per iterazione a questi livelli non è eccessivamente elevato.

Per di più, ad un fissato livello l , avendo bisogno di calcolare $\|\phi_2^l - \phi_0^l\|_1$, sarà possibile applicare il test (4.19) soltanto a partire da $m = 2$, supponendo sempre che il valore iniziale di m sia 0. Dunque forzeremo ad ogni livello che il solutore iteri il processo per lo meno fino alla seconda iterazione (fino al calcolo di ϕ_m^l con $m = 2$).

Considerando (4.19), decidiamo di arrestare il programma ad un livello l tale che l'errore ϵ^l sia minore della tolleranza tol data. Il test di arresto implementato nel solutore sarà

$$\eta^l < tol. \quad (4.21)$$

Prima di passare alla descrizione dei programmi, dobbiamo discutere su come implementare il test (4.19). Non potendo calcolare $\|\{\phi(x_i^l)\}\|_1$, lo approssimeremo con la sua miglior stima in norma $\|\cdot\|$ che abbiamo al momento in cui è necessario effettuare il test. Ci rimane solo da considerare l'approssimazione di $\|\tilde{\phi}_m^l - \{\phi(x_i^l)\}\|_1$. Nel caso del solutore che si basa sul metodo di Broyden, usiamo lo step $\|phi_{m-1}^l - \phi_m^l\|_1$, avendo visto che si tratta di una buona approssimazione per i metodi di tipo Newton in generale [11].

Quando invece il solutore è basato su un metodo di tipo Picard usiamo la stima

$$\|\tilde{\phi}_m^l - \{\phi(x_i^l)\}\|_1 \approx \frac{\theta}{1 - \theta} \|\phi_m^l - \phi_{m-1}^l\|_1, \quad (4.22)$$

dove θ è la costante asintotica [6]. Il solutore approssima θ al livello 0 con la media dei $\|\phi_{k+2}^0 - \phi_{k+1}^0\|_1 / \|\phi_{h+1}^0 - \phi_k^0\|_1$ con $0 \leq k \leq m-2$, dove m^l indica il massimo numero di iterazioni ad un livello l . L'approssimazione di θ viene aggiornata nei livelli successivi se $m^l \geq 3$, scelta giustificata dal fatto che se $m^l = 1, 2$ l'approssimazione di θ potrebbe non essere attendibile. Abbiamo notato sperimentalmente che θ da noi così approssimato tende ad essere indipendente dal livello.

Veniamo dunque alla descrizione dei programmi. Scriviamo uno pseudocodice per un solutore che al suo interno chiamerà il sottoprogramma *sol*, che sarà *brsol* nel caso in cui il solutore si basi sul metodo di Broyden, *uppic* o *pic* quando si basa rispettivamente su Updated Picard o Picard semplice.

Procedura 4.0.1 $\text{ml}(N, c, \text{intfin}, \alpha, \beta, \text{maxit}, \text{nmax}, \text{tol}, \phi, K, f, a, b)$

- (1.) calcolo dei pesi e nodi di Fejer al livello 0
- (2.) for $i = 0, N$
 $u_0^0(i) = 0$
- (3.) call **sol**($u^0, x^0, w^0, u^1, x^1, w^1, ti, tiv, \text{maxit}, \text{nmax}, l, \text{tol}, K, f, a, b$)
- (4.) repeat
 - (i) $N^{l+1} = N^l * (1 + c)$, $l = l + 1$
 - (ii) calcolo dei pesi e nodi di Fejer al livello 1
 - (iii) call **sol**($u^{l-1}, x^{l-1}, w^{l-1}, u^l, x^l, w^l, ti, tiv, \text{maxit}, \text{nmax}, l, \text{tol}, K, f, a, b$)
 - until $(\|\phi_2^l - \phi_0^l\|_1^2) / (\|\phi_2^{l-1} - \phi_0^{l-1}\|_1 \cdot \|\phi_m^l\|_*) < \text{tol}$
- (5.) end

Scriviamo ora il sottoprogramma *brsol* che implementa il metodo risolutivo di Broyden.

brsol($u^{l-1}, x^{l-1}, w^{l-1}, u^l, x^l, w^l, ti, tiv, \text{maxit}, \text{nmax}, l, \text{tol}, K, f, a, b$)

- (1.) if $l > 0$ then
 interpolazione di Nystrom per u^{l-1} dal livello $l-1$ a l
- (2.) for $i = 0, N$ $j = 0, N$
 $M(i, j) = w(j)k(x^l(i), x^l(j))$

- (3.) $s_0 = -Fu_0^l$
- (4.) $u_1 = u_0^l + s_0$
- (5.) $s_1 = Fu_1^l / (1 + s_0^T * Fu_1^l / \|s_0\|_2)$
- (6.) $m = 1$ itc = 2
- (7.) if (itc > maxit) goto (15.)
- (8.) $m = m + 1$ itc = itc + 1
- (9.) $u_m^l = u_{m-1}^l + s_{m-1}$
- (10.) if $m = 2$ calcolare $\|\phi_2^l - \phi_0^l\|_1$
- (11.) if $l \geq 2$ then
 $tol_i = \max\{\|\phi_2^l - \phi_0^l\|_1^2 / (\|\phi_2^{l-1} - \phi_0^{l-1}\|_1 \|\phi_m^l\|_1), tol\}$
else
 $tol_i = tol$
- (11.) if $\|\phi_m^l - \phi_{m-1}^l\|_1 / \|\phi_m^l\|_1 \leq tol_i$ goto (15.)
- (12.) if $m < nmax$ then
 - (i) $z = -Fu_m^l$
 - (ii) for $j = 1, m - 1$
 $z = z + s_j s_{j-1}^T z / \|s_{j-1}\|_2^2$
 - (iii) $s_m = z / (1 + s_{m-1}^T z / \|s_{m-1}\|_2^2)$
- (13.) if $m = nmax$ then
ripetere i punti (c), (d), (e)
 $m = 1$
- (14.) goto (7.)
- (15.) continue
return

La convergenza del metodo di Broyden è localmente superlineare, ed essendo un ottimo metodo risolutivo per equazioni non lineari, si presenta come un valido termine di confronto per i metodi di tipo Picard.

Ora riportiamo lo pseudo-codice per il sottoprogramma che utilizza lo schema iterativo *Updated Picard*.

`uppic( $u^{l-1}, x^{l-1}, w^{l-1}, u^l, x^l, w^l, ti, tiv, maxit, l, tol, k, f, a, b$ )`

- (1.) if $l > 0$ then
interpolazione di Nystrom per u^{l-1} dal livello $l - 1$ a l
- (2.) for $i = 0, N$ $j = 0, N$
 $M(i, j) = w(j)k(x^l(i), x^l(j))$
- (3.) $m = 0$
- (4.) repeat
 - (i) $m = m + 1$
 - (ii) for $i = 0, N$
 $z(i) = \sum_{j < i} M(i, j)f(x^l(j), z(j)) + \sum_{j \geq i} M(i, j)f(x^l(j), u_{m-1}^l(i))$
 - (iii) $u_m^l = (1 - \omega)u_{m-1}^l + \omega z$
 - (iv) if $m \geq 2$ then
 - if $l \geq 2$ then
 $tol_i = \max\{\|\phi_2^l - \phi_0^l\|_1^2 / (\|\phi_2^{l-1} - \phi_0^{l-1}\|_1 \|\phi_m^l\|_1), tol\}$
 - else
 $tol_i = tol$
 - if $l = 0$ then
 $rerr = \|\phi_{m-1}^l - \phi_m^l\|_1 / \|\phi_m^l\|_1$
 - else
 $rerr = (\|\phi_{m-1}^l - \phi_m^l\|_1 \theta) / (\|\phi_m^l\|_1 (1 - \theta))$
until ($rerr < tol_i$ and $m \geq 2$)
- (5.) if $m \geq 3$ then
 $\theta = \left(\sum_{i=2}^m \|\phi_i^l - \phi_{i-1}^l\|_1 / \|\phi_{i-2}^l - \phi_{i-1}^l\|_1 \right) / (m - 1)$
- (6.) return

Infine lo pseudo-codice per il sottoprogramma che utilizza lo schema iterativo *Picard*. Sottolineiamo che il sottoprogramma *pic* è identico ad *uppic* tranne che per il punto (4.)(ii).

pic($u^{l-1}, x^{l-1}, w^{l-1}, u^l, x^l, w^l, ti, tiv, maxit, l, tol, k, f, a, b$)

- (1.) if $l > 0$ then
interpolazione di Nystrom per u^{l-1} dal livello $l - 1$ a l
- (2.) for $i = 0, N$ j = 0, N
 $M(i, j) = w(j)k(x^l(i), x^l(j))$
- (3.) $m = 0$
- (4.) repeat
 - (i) $m = m + 1$
 - (ii) for $i = 0, N$
 $z(i) = \sum_{j=0}^{N^l} M(i, j)f(x^l(j), u_{m-1}^l(i))$
 - (iii) $u_m^l = (1 - \omega)u_{m-1}^l + \omega z$
 - (iv) if $m \geq 2$ then
 - if $l \geq 2$ then
 $toli = \max\{\|\phi_2^l - \phi_0^l\|_1^2 / (\|\phi_2^{l-1} - \phi_0^{l-1}\|_1 \|\phi_m^l\|_1), tol\}$
 - else
 $toli = tol$
 - if $l = 0$ then
 $rerr = \|\phi_{m-1}^l - \phi_m^l\|_1 / \|\phi_m^l\|_1$
 - else
 $rerr = (\|\phi_{m-1}^l - \phi_m^l\|_1 \theta) / (\|\phi_m^l\|_1 (1 - \theta))$
until ($rerr < toli$ and $m \geq 2$)
- (5.) if $m \geq 3$ then
 $\theta = \left(\sum_{i=2}^m \|\phi_i^l - \phi_{i-1}^l\|_1 / \|\phi_{i-2}^l - \phi_{i-1}^l\|_1 \right) / (m - 1)$
- (6.) return

La convergenza dei due ultimi metodi è garantita dai teoremi visti nel capitolo 2 per il caso finito-dimensionale.

4.1 Esempi numerici.

Riportiamo ora i risultati dei test effettuati sui solutori NFB (solutore con metodi di Broyden, formula interpolatoria di Nyström e nodi e pesi di Fejer), PNF_ω (metodi di Picard semplice con costante di rilassamento w , formula interpolatoria di Nyström e nodi e pesi di Fejer) e $UPNF_\omega$ (metodo di Updated Picard con costante di rilassamento ω , formula interpolatoria di Nyström e nodi e pesi di Fejer). Confronteremo i solutori sulla risoluzione dell'equazione (4.1) in vari casi sia con nucleo di Boltzmann che con nucleo di Chandrasekhar. Abbiamo calcolato le costanti di rilassamento ottimali con una semplice procedura di bisezione. In ogni tabella, oltre al tipo di problema, riportiamo il numero di nodi per livello, la tolleranza relativa, l' ω ottimo e il numero di iterazioni per ogni programma ad ogni livello.

Le prime 6 tabelle prendono in esame il caso del nucleo di Boltzmann. La tolleranza relativa è 10^{-6} . Si nota come le prime 4 combinazioni di p e q generano delle equazioni più “facili” rispetto a quelle riportate nelle tabelle 1.5 e 1.6, se si considera il numero di livelli necessari per arrivare alla precisione 10^{-6} . —E' comune a tutti i casi di nucleo di Boltzmann che al crescere di λ , cresce la “difficoltà” del problema, perché pur restando invariato il numero di livelli, cresce il numero di iterazioni per livelli. È interessante notare che al variare delle combinazioni di p e q nelle prime 6 tabelle, ω ottimo si discosta di poco a parità di λ nei singoli solutori.

Tabella 4.1: Nucleo di Boltzmann, $p = 0$, $q = 0$, $tol = 10^{-6}$.

	# di punti	8	16	32
$\lambda = 1$	$NFP_{0.80}$	4	2	2
	$NFUP_{0.87}$	6	4	2
	$NFBroyden$	5	4	3
$\lambda = 10^2$	$NFP_{0.55}$	8	4	2
	$NFUP_{0.73}$	10	6	3
	$NFBroyden$	10	4	3
$\lambda = 10^4$	$NFP_{0.51}$	10	4	2
	$NFUP_{0.73}$	12	8	3
	$NFBroyden$	14	4	3

In generale per il nucleo di Boltzmann si vede che i metodi di tipo Picard sono per numero di iterazioni per livello paragonabili se non migliori del me-

Tabella 4.2: Nucleo di Boltzmann, $p = 0$, $q = 1$, $tol = 10^{-6}$.

# di punti		8	16	32
$\lambda = 1$	$NFP_{0.77}$	7	5	2
	$NFUP_{0.84}$	7	4	2
	NFB	6	5	4
$\lambda = 10^2$	$NFP_{0.68}$	13	9	4
	$NFUP_{0.73}$	10	6	2
	NFB	10	6	5
$\lambda = 10^4$	$NFP_{0.64}$	16	11	4
	$NFUP_{0.73}$	12	7	2
	NFB	15	6	2

todo di Broyden, riscontrando un'eccellente prestazione di Picard semplice, cosa che non succede per l'equazione di Chandrasekhar (Tabella 1.6).

Da questa tabella risulta che il problema diventa più dispendioso quando λ si avvicina ad $1/2$. Questa maggiore difficoltà però non fa aumentare il numero di livelli, ma solo il numero di iterazioni per livello. Sempre prendendo in considerazione il numero di iterazioni per livello si vede che i metodi migliori sono Broyden e Updated Picard, mentre Picard semplice ha bisogno di un numero di iterazioni maggiori per arrivare alla stessa precisione.

Tabella 4.3: Nucleo di Boltzmann, $p = 1$, $q = 0$, $tol = 10^{-6}$.

# di punti		8	16	32
$\lambda = 1$	$NFP_{0.77}$	5	2	2
	$NFUP_{0.88}$	6	4	2
	NFB	6	4	3
$\lambda = 10^2$	$NFP_{0.55}$	8	4	2
	$NFUP_{0.79}$	9	5	3
	NFB	10	4	3
$\lambda = 10^4$	$NFP_{0.51}$	11	4	2
	$NFUP_{0.73}$	12	8	3
	NFB	15	4	3

Tabella 4.4: Nucleo di Boltzmann, $p = 1$, $q = 1$, $tol = 10^{-6}$.

# di punti		8	16	32
$\lambda = 1$	$NFP_{0.82}$	8	5	2
	$NFUP_{0.91}$	6	3	2
	NFB	7	5	3
$\lambda = 10^2$	$NFP_{0.69}$	14	9	4
	$NFUP_{0.86}$	9	5	3
	NFB	12	6	5
$\lambda = 10^4$	$NFP_{0.68}$	16	12	5
	$NFUP_{0.88}$	12	7	3
	NFB	15	6	5

Tabella 4.5: Nucleo di Boltzmann, $p = 1$, $q = -1$, $tol = 10^{-6}$.

	# di punti	64	128	256	512	1024	2048
$\lambda = 1$	$NFP_{0.82}$	7	3	2	2	2	2
	$NFUP_{0.95}$	6	2	2	2	2	2
	NFB	7	4	2	2	2	2
$\lambda = 10^2$	$NFP_{0.69}$	11	6	2	2	2	2
	$NFUP_{0.91}$	8	4	2	2	2	2
	NFB	11	5	3	3	3	3
$\lambda = 10^4$	$NFP_{0.66}$	15	8	2	3	2	2
	$NFUP_{0.87}$	10	4	2	2	2	2
	NFB	15	5	3	3	3	3

Tabella 4.6: Nucleo di Boltzmann, $p = 2$, $q = -1$, $tol = 10^{-6}$.

	# di punti	64	128	256	512	1024	2048
$\lambda = 1$	$NFP_{0.82}$	8	4	2	2	2	
	$NFUP_{0.91}$	7	3	2	2	2	
	NFB	9	4	3	3	2	
$\lambda = 10^2$	$NFP_{0.68}$	13	6	2	2	2	
	$NFUP_{0.91}$	8	4	2	2	2	2
	NFB	13	5	3	3	3	2
$\lambda = 10^4$	$NFP_{0.64}$	16	7	2	2	2	2
	$NFUP_{0.91}$	10	5	2	2	2	2
	NFB	16	5	3	3	3	3

Tabella 4.7: Equazione di Chandrasekhar, $tol = 10^{-8}$.

	# di punti	8	16	32	64
$\lambda = 0.3$	$NFP_{0.86}$	8	4	2	2
	$NFUP_{0.96}$	7	3	2	2
	NFB	7	5	2	2
$\lambda = 0.45$	$NFP_{0.76}$	11	5	3	2
	$NFUP_{0.91}$	9	4	2	2
	NFB	7	5	2	2
$\lambda = 0.495$	$NFP_{0.67}$	15	7	3	2
	$NFUP_{0.87}$	10	4	2	2
	NFB	11	6	2	2
$\lambda = 0.49995$	$NFP_{0.64}$	18	9	3	2
	$NFUP_{0.82}$	13	5	2	2
	NFB	15	6	2	2
$\lambda = 0.4999995$	$NFP_{0.64}$	20	10	3	2
	$NFUP_{0.86}$	14	6	2	2
	NFB	19	6	2	2

Bibliografia

- [1] Atkinson, K.E.: A survey of numerical methods for solving nonlinear integral equations, *J. Integral Equations Appl.* **4** (1992), 15-46.
- [2] Boffi, V.C. and Spiga, G.: An equation of Hammerstein type arising in particle transport theory, *J. Math. Phys.* **24** (1983), 1625-1629.
- [3] Boffi, V.C., Spiga, G. and Thomas J.R.: Solution of nonlinear integral equation arising in particle transport theory, *J. Comput. Phys.* **59** (1985), 96-107.
- [4] Cercignani, C.: *The Boltzmann Equation and Its Applications*, Springer, New York, 1988.
- [5] Chandrasekhar, S.: *Radiative Transfer*, Dover, New York, 1960.
- [6] Comincioli
- [7] Gautschi, W.: Orthogonal polynomials: applications and computations, *Acta Numerica* (1996), 45-119.
- [8] Guo, D. e Lakshmikantham, V.: *Nonlinear Problems in Abstract Cones*, Academic Press Inc., New York, 1988.
- [9] Hively, G.A.: On a class of nonlinear integral equations arising in transport theory, *SIAM J. Math. Anal.* **9** (1978), 787-792.
- [10] Hu, S., Khavanin, M., e Zhuang, W.: Integral equations arising in the kinetic theory of gases, *Appl. Anal.* **34** (1989), 261-266.
- [11] Kelley, C.T.: *Iterative Methods for Linear and Nonlinear Equations*, Frontiers in Applied Mathematics 16, SIAM, Philadelphia, 1995.

- [12] Sommariva A.: *Constructive and numerical analysis for a class of Hammerstein equations arising in transport theory*, Doctoral Dissertation in Computational Mathematics, University of Padova, 1998, advisors: Prof. I. Moret and Dr. M. Vianello.
- [13] Sommariva A. and Vianello M.: Computing fixed-points of decreasing operators by relaxed iterations, spedito per la pubblicazione.
- [14] Sommariva A. and Vianello M.: Constructive analysis of purely integral Boltzmann models, *J. Integral Equations Appl.*, in corso di stampa.
- [15] Sommariva A. and Vianello M.: Approximating fixed-points of decreasing operators in spaces of continuous functions, *Numer. Funct. Anal. and Optim.* **19** (1998), 635-646.
- [16] Sommariva A., Vianello M. and Facchinello E.: Relaxed Picard-like methods for nonlinear integral equations arising in transport theory.