



UNIVERSITÀ DEGLI STUDI DI PADOVA

Dipartimento di Fisica e Astronomia Galileo Galilei

Corso di Laurea magistrale in Astronomia

Identificazione di buchi neri di massa intermedia attraverso tecniche di machine learning interpretabile

Relatore
Prof.ssa Michela Mapelli

Laureanda
Erica Greco

Correlatore
Dott. Mario Pasquato

24 Settembre 2020
Anno Accademico 2019-2020

Un pretesto per tornare bisogna sempre
seminarselo dietro, quando si parte.
A. Baricco, Oceano mare

Abstract

Vari scenari teorici indicano che gli ammassi globulari possano ospitare buchi neri di massa intermedia (*Intermediate-Mass Black Holes*, IMBHs), ma gli approcci osservativi tentati finora non ne hanno ancora dimostrato la presenza in modo incontrovertibile. Una rilevazione diretta per mezzo di emissione elettromagnetica non è infatti ancora stata effettuata, in parte a causa dell'ambiente privo di gas che caratterizza gli ammassi globulari. Per questo motivo sono stati sviluppati metodi di identificazione indiretta.

Recentemente, approcci indiretti basati sulla cinematica delle pulsar quale quello proposto da [3] hanno riscosso notevole interesse. Nel lavoro di tesi si è utilizzata una tecnica innovativa per prevedere la presenza di un IMBH per mezzo delle pulsar, applicandola per ora a dati simulati. Nello specifico, ci si riferisce all'utilizzo di un modello di *Machine Learning* (ML) interpretabile: gli alberi decisionali di classificazione.

Per costruire il *dataset* su cui allenare il modello ci si è serviti di un campione di ammassi simulati con un codice diretto a N-corpi, in cui alcune particelle situate nelle regioni centrali sono state etichettate come pulsar, ottenendone una serie di misure intese a simulare possibili osservabili, quali la posizione proiettata, la velocità lungo la linea di vista, l'accelerazione e le sue derivate. L'albero decisionale appreso dal modello nel caso in cui si utilizzino 20 pulsar è in grado di classificare gli ammassi secondo la presenza di un IMBH al loro interno sulla base di tali

caratteristiche con il 70% di accuratezza. Tale risultato presenta già una possibile utilità per un'applicazione a dati reali finalizzata alla selezione preliminare di ammassi candidati ad uno studio più dettagliato con altri mezzi. Inoltre, per costruzione, l'albero decisionale si presta ad un'analisi delle motivazioni per cui il modello ha effettuato una data predizione (interpretabilità), che cerco di spiegare in termini di interpretazione fisica in questa tesi.

Indice

Introduzione	1
1 A caccia di buchi neri di massa intermedia	4
1.1 I buchi neri di massa intermedia	5
1.2 Metodi di identificazione di IMBHs al centro di ammassi globulari	9
1.3 Metodo delle MSPs	12
1.3.1 Profili radiali di Jerk e Snap in ammasso	13
1.3.2 Metodo osservativo	15
1.3.3 Simulazioni e sviluppi con modelli di Machine Learning	17
2 Tecniche di Machine Learning interpretabile: alberi deci- sionali	19
2.1 Introduzione al Machine Learning	19
2.2 Concetti preliminari: notazione e struttura del dataset . .	22
2.3 Alberi decisionali: concetti generali	23
2.4 Algoritmo CART	26
2.5 Criteri per la costruzione di alberi decisionali	27
2.6 Regole di Splitting	28
2.6.1 La funzione di impurità	29
2.7 Dichiarazione dei nodi terminali	30

2.7.1	Criteri di arresto	31
2.7.2	Assegnazione della classe al nodo terminale	32
2.8	Problema dell'overfitting	32
2.9	Cost-Complexity Pruning	34
2.9.1	Notazione e metodo	35
2.10	Training, validation e test sets	38
2.11	Qualità delle previsioni	39
2.11.1	Accuratezza degli alberi decisionali	39
2.11.2	Matrice di confusione: precisione, richiamo e F-score	40
3	Alberi decisionali per l'identificazione di IMBHs	44
3.1	Motivazioni per l'uso degli alberi decisionali	45
3.2	Simulazioni con il codice N-body diretto HiGPUs	46
3.3	Simulazioni e condizioni iniziali	49
3.4	Dataset	52
3.5	Codice	54
4	Risultati	57
4.1	Valutazione delle performance	57
4.2	Interprtazione dei risultati	62
5	Conclusioni e sviluppi futuri	69
	Riferimenti bibliografici	72

Introduzione

I buchi neri di massa intermedia (*Intermediate-Mass Black Holes*, IMBHs) sono buchi neri la cui massa è compresa tra 10^2 e $10^6 M_\odot$. Si pensa che essi possano essere l'anello mancante necessario per collegare le informazioni note riguardanti i buchi neri di origine stellare ($\sim 10 M_\odot$) a quelle riguardanti i buchi neri super-massicci ($10^6 - 10^{10} M_\odot$). Infatti, i modelli teorici prevedono l'esistenza degli IMBHs al centro di ammassi globulari (*Globular Cluster*, GC), mentre le osservazioni condotte finora non ne hanno ancora rilevato la presenza, fatta eccezione del recentissimo evento di detection di onde gravitazionali GW190521, che viene interpretato come la formazione di un buco nero di $\approx 150 M_\odot$ [4], quindi in ogni caso di massa abbastanza vicina all'estremo stellare del range di massa degli IMBH.

La motivazione per cui questi oggetti peculiari risultano sfuggenti rispetto ai buchi neri appartenenti alle altre categorie è legata principalmente al fatto che la strumentazione attuale presenti dei limiti in tal senso. Tuttavia, esistono dei metodi indiretti per poter identificare gli IMBHs negli ammassi. Uno fra essi è quello che si basa sulla dinamica delle MSPs nei GCs. In particolare, attraverso misure delle derivate dei periodi delle MSPs sarebbe possibile ottenere i loro valori di accelerazione e delle sue derivate: il *jerk* (derivata prima dell'accelerazione) e lo *snap* o *jounce* (derivata seconda dell'accelerazione). Dallo studio di tali caratteristiche, poi, si potrebbe risalire all'identificazione degli IMBHs

[3].

A partire da questa proposta, questo lavoro si sviluppa verso tale scopo, ma utilizzando metodi diversi. Infatti, vengono utilizzate tecniche di *Machine Learning* interpretabile per prevedere la presenza di IMBHs al centro di GCs.

Poiché i dati osservativi in questo campo non sono sufficienti, per lo studio si sono utilizzati dati ottenuti da un set di simulazioni a N-corpi. Il modello di *Machine Learning* utilizzato è quello degli alberi decisionali, caratterizzati dalla possibilità di essere interpretati. Questa è un'importante caratteristica per un modello previsionale che riguarda soprattutto la fase di analisi dei risultati.

Gli ammassi simulati sono stati classificati dal modello secondo la presenza di un IMBH o meno sulla base delle caratteristiche dinamiche delle MSPs scelte nelle regioni centrali dei GCs. Nello specifico, le caratteristiche fornite al modello sono: la distanza delle MSPs dal centro di massa del GC proiettata sul piano del cielo e le componenti di velocità, accelerazione, *jerk* e *snap*, considerate lungo la linea di vista.

La ricerca di IMBHs attraverso questo metodo potrebbe essere una nuova strada alternativa da percorrere per ottenere dei risultati soddisfacenti, specie se tali metodi in futuro potranno essere applicati a dati reali.

Il lavoro di tesi può essere suddiviso principalmente in tre parti. La prima e la seconda parte si sviluppano rispettivamente nei capitoli 1 e 2. La terza parte, invece, contiene i restanti tre capitoli.

Nel primo capitolo viene presentato e contestualizzato il problema astrofisico su cui si basa l'intera tesi: la ricerca degli IMBHs in ammassi globulari. Viene, dunque, fornita un'introduzione sugli IMBHs e su quali sono i metodi utilizzati finora per identificarli, per poi concentrarsi, in particolare, sul metodo delle MSPs.

Nel secondo capitolo, invece, viene descritto lo strumento utilizzato per lo sviluppo del progetto, ovvero gli alberi decisionali. Dopo una breve

introduzione al *Machine Learning*, viene descritto l'algoritmo CART [14] che è quello utilizzato per la costruzione degli alberi decisionali in questo lavoro. Vengono, infatti, descritte le fasi in cui l'algoritmo si sviluppa e le metodologie che esso adopera, come, ad esempio, la tecnica del *pruning* utilizzata in questo caso per risolvere un problema comune a tutti gli algoritmi di *Machine Learning*: il problema dell'*overfitting*. Inoltre, in questo capitolo, vengono forniti gli strumenti per valutare le prestazioni di un algoritmo di classificazione.

Infine, nell'ultima parte si entra nel merito del lavoro di tesi descrivendo nello specifico i *dataset* utilizzati, il codice sviluppato per la classificazione degli ammassi (Cap. 3) e, in ultimo, vengono presentati i risultati ottenuti con la relativa interpretazione (Cap. 4). Nel quinto ed ultimo capitolo, invece, vengono discusse le conclusioni e presentate le possibili prospettive future per l'identificazione di IMBHs in GCs attraverso tecniche di *Machine Learning*.

Capitolo 1

A caccia di buchi neri di massa intermedia

In questo Capitolo viene presentato il *background* astrofisico sul quale si sviluppa l'intero lavoro di tesi.

Ci si focalizza sull'importanza di voler ricercare i buchi neri di massa intermedia (*Intermediate-Mass Black Holes*, IMBHs) all'interno degli ammassi globulari (*Globular Clusters*, GCs) e vengono presentate alcune tecniche utilizzate per la loro identificazione. In particolare, viene approfondito il metodo di identificazione basato sulla dinamica delle Pulsar Millisecondo (*Millisecond Pulsars*, *MSPs*) all'interno dei GCs. Infatti, questo progetto si sviluppa in una direzione diversa, ma complementare all'idea suggerita da Abbate et al.(2019) [3] di identificare gli IMBHs nei GCs utilizzando, oltre alle velocità e accelerazioni delle MSPs in ammasso, anche le informazioni dinamiche fornite dalle derivate dell'accelerazione.

Nello specifico, nella mia tesi, ho sviluppato un codice che utilizza un modello di *Machine Learning* (ML) interpretabile in grado di fare previsioni sulla presenza di IMBHs in ammassi sulla base di tali informazioni

dinamiche delle MSPs.

1.1 I buchi neri di massa intermedia

I buchi neri vengono classificati in base alla loro massa in: buchi neri di origine stellare, buchi neri di massa intermedia e buchi neri super-massicci.

I buchi neri di origine stellare sono oggetti compatti che si formano alla fine della vita di stelle massicce in seguito ad esplosioni di Supernova, mentre i buchi neri super-massicci normalmente si trovano al centro di galassie. I buchi neri di massa intermedia, invece, (*Intermediate-Mass Black Holes*, IMBHs) sono buchi neri la cui massa (tra 100 e $10^6 M_{\odot}$) è significativamente maggiore di quella dei buchi neri di origine stellare e al contempo molto minore di quella dei buchi neri super-massicci.

I buchi neri sono caratterizzati da intensi campi gravitazionali dai quali, escludendo fenomeni quantistici non ancora verificati sperimentalmente, nulla può fuoriuscire, nemmeno la radiazione. Per questo motivo non possono essere osservati direttamente, ma possiamo unicamente misurare gli effetti che essi hanno sugli altri corpi celesti. Ad esempio, è possibile rilevare le emissioni X quando essi, in un sistema binario, accrescono sottraendo massa ad una stella compagna (Fig.1.1) o le emissioni sottoforma di onde gravitazionali quando si fondono (fenomeno di *merging*) con altri oggetti compatti o anche si possono rilevare gli effetti gravitazionali che hanno sulle orbite di altri oggetti, come nel caso di un buco nero super-massivo nel centro galattico.

Gli IMBHs, oggetto di studio di questa tesi, dal punto di vista osservativo, potrebbero rappresentare il nesso che collega le informazioni raccolte finora e che riguardano principalmente i buchi neri di origine

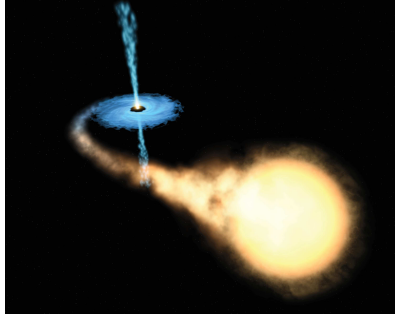


Figura 1.1: Esempio, in un *artist impression*, di buco nero in un sistema binario che accresce la propria massa a discapito della sua stella compagna [46].

stellare e quelli super-massici. Infatti, nonostante i modelli teorici suggeriscano fortemente che è possibile trovare gli IMBHs all'interno di ammassi globulari [34], attualmente non si dispone di prove definitive per confermarlo, in quanto gli indizi che abbiamo sono ancora relativamente controversi e non si è ancora raggiunto il consenso presso la comunità scientifica (Fig.1.2).



Figura 1.2: Vuoto di informazioni tra le categoria di buchi neri di origine stellare ed i buchi neri super-massivi: gli IMBHs.

Gli scenari di formazione proposti per questi oggetti peculiari sono molteplici [22].

Il primo scenario ipotizza che gli IMBHs possono essersi originati in seguito al collasso di stelle di Popolazione III, ovvero ipotetiche stelle formatesi

con abbondanza chimica primordiale che possono raggiungere alte masse rispetto alle stelle odierne [5].

La seconda ipotesi è quella del collasso diretto secondo cui la loro massa si sarebbe addensata direttamente dal materiale dell'universo in formazione subito dopo il *Big Bang*, senza passare per le fasi dell'evoluzione stellare [20].

Un ultimo interessante scenario, invece, suggerisce che essi si formino mediante fenomeni di collisione e di *merging*. In questo caso esistono più ipotesi.

Secondo *Miller & Hamilton 2002* [34] un buco nero di origine stellare di $50 M_{\odot}$ potrebbe collocarsi in breve tempo al centro dell'ammasso per l'effetto di segregazione di massa; successivamente, collisioni con altri buchi neri stellari mediate da incontri gravitazionali a tre e quattro corpi e da perdita di energia per emissione di onde gravitazionali, porterebbero al raggiungimento di masse dell'ordine di $10^3 M_{\odot}$ in un tempo paragonabile a quello di Hubble.

Un altro meccanismo, invece, propone che in un *core* ad alta densità, le stelle molto massicce ($50 - 100 M_{\odot}$) possono essere soggette ad un'efficiente segregazione di massa che le colloca nel nucleo dell'ammasso mentre si trovano in fase di Sequenza Principale. A questo punto si verificherebbe un numero sempre crescente di collisioni e *merging* tra stelle che porterebbero alla formazione e all'immediato collasso di una stella con massa pari a $\simeq 10^{-3}$ volte la massa dell'ammasso, generando così un IMBH [40].

I GCs, essendo ambienti stellari densi e dinamicamente attivi, potrebbero essere i luoghi ideali per la formazione di IMBHs attraverso collisioni stellari o incontri gravitazionali tra buchi neri di origine stellare e successive fusioni [41]. L'interesse per la formazione di IMBHs in GCs è anche legato al fatto che gli ammassi globulari più massicci tendono a precipitare rapidamente verso il centro della galassia, dove potrebbero concorrere a

formare il *nuclear star cluster* [10]. Pertanto, essi rappresentano un possibile meccanismo di formazione di un buco nero super-massivo attraverso fenomeni di *merging* tra IMBHs.

In condizioni di prossimità ad una binaria di oggetti compatti gli IMBHs possono essere sorgenti di emissione di onde gravitazionali.

Gli interferometri terrestri di onde gravitazionali LIGO e Virgo [26], [48], [6] operano in un intervallo di frequenze tra qualche decina e qualche migliaio di Hertz. Quindi possono osservare la coalescenza di buchi neri fino ad una massa di $\sim 500M_{\odot}$, nel regime degli IMBHs.

La prossima generazione di interferometri nello spazio, come il Laser Interferometer Space Antenna, LISA [7], sarà ancora più adatta ad osservare gli IMBHs, poiché opererà nel range di frequenza compreso tra $10^{-3} - 10^{-1}$ Hertz.

Inoltre, gli IMBHs sono molto difficili da individuare in quanto, gli ammassi globulari sono ambienti decisamente poveri di gas, poiché sono costituiti principalmente da stelle vecchie e il gas primordiale che era presente è stato utilizzato nelle precedenti fasi di formazione stellare o spazzato via da esplosioni di Supernova. Per tale ragione e per come operano i meccanismi di accrescimento, ovvero tramite la dissipazione di energia e momento angolare verso l'esterno di un disco, risulta difficile rilevare emissioni X, determinate da tali fenomeni.

Un'altra ragione per cui risulta difficoltoso rilevarli è legata alla componente stellare.

Per ogni buco nero di massa M_{BH} , infatti, è possibile stimare il raggio di influenza r_i :

$$r_i = \frac{GM_{BH}}{\sigma^2} \quad (1.1)$$

con G la costante di gravitazione universale e σ la dispersione di velocità delle stelle subito al di fuori della sfera di influenza del buco nero. Le stelle all'interno di questo raggio risentono principalmente dell'influenza gravitazionale del buco nero stesso.

Il raggio di influenza di un IMBH è dell'ordine di solo qualche secondo d'arco. Ad esempio, un IMBH di $\sim 1000M_{\odot}$ all'interno di *47 Tucanae* avrebbe un raggio di influenza di circa 1" [22], contro la dimensione apparente dell'ammasso di circa 31' [1]. Pertanto, in generale, la sfera di influenza di un IMBH non è così importante, rispetto alle dimensioni dell'ammasso, da influenzare gravitazionalmente un numero di stelle sufficientemente elevato da determinare una quantità di eventi di accrescimento mareale tale da poter essere osservato.

1.2 Metodi di identificazione di IMBHs al centro di ammassi globulari

Di fronte alla presenza di un IMBH in un ammasso stellare, ci aspettiamo che si verifichino delle alterazioni nei profili di densità del GC e di dispersione di velocità delle stelle dell'ammasso. Sulla base di questi si costruiscono la maggior parte dei metodi di identificazione degli IMBHs con le relative difficoltà a livello pratico. Inoltre, se ci trovassimo in presenza di un fenomeno di accrescimento, la misurazione delle emissioni X e radio sarebbe un altro canale identificativo di IMBH al centro di un GC [47].

Un IMBH aumenta la profondità della buca di potenziale di un ammasso causando un incremento della densità stellare nelle zone centrali. Queste ultime, pertanto, risultano avere un profilo di densità di una *cuspid*e ben descritta dalla legge di potenza $\rho \propto r^{-\alpha}$, con $\alpha \simeq 1.55$ [13].

Le stelle all'interno del raggio di influenza r_i (eq. 1.1) del IMBH, oltre a risentire principalmente dell'influenza gravitazionale del buco nero, seguono un profilo di dispersione di velocità di tipo *kepleriano* caratterizzato da una ripidità nelle regioni centrali. Tuttavia, la determinazione accurata della dispersione di velocità delle stelle valutata solo nella regione centrale del GC, è al limite delle capacità dell'attuale strumentazione

astronomica e richiede un'elevatissima risoluzione spaziale. I risultati di misurazioni effettuate con metodi diversi (spettroscopia in luce integrata, stelle individuali, studio dei moti propri) sono talvolta in disaccordo tra loro, con la conseguenza che la comunità scientifica non ha ancora raggiunto un consenso in merito [30].

Il profilo di dispersione di velocità, però si può ottenere teoricamente sulla base del profilo di densità determinato precedentemente e considerando un oggetto puntiforme di una certa massa al centro dell'ammasso [29].

Come risultato si ottiene una famiglia di profili di dispersione di velocità che, una volta confrontati con le osservazioni, restituiscono la massa del BH centrale (Fig.1.3).

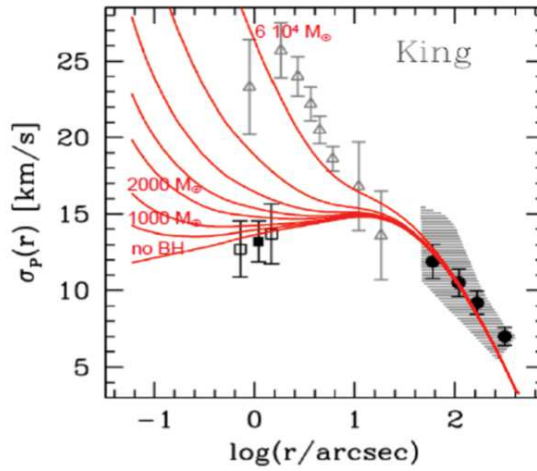


Figura 1.3: Profilo di dispersione di velocità osservato per l'ammasso globulare NGC 6388 con sovrapposte le linee rosse continue delle famiglie di modelli considerate. Esse sono state ottenute assumendo una diversa massa per il buco nero centrale. Si noti che aumentando la massa del IMBH la pendenza nelle regioni centrali risulta più ripida [25], [29].

Tuttavia, la presenza di una *cuspide* nelle regioni centrali di un GC

può essere dovuta anche a cause diverse dalla presenza di un IMBH. Ad esempio, a seguito di un collasso del *core*, tanti buchi neri di massa stellare potrebbero trovarsi nelle regioni centrali dell'ammasso e determinare, quindi, i ripidi profili di velocità e densità associati alle cuspidi [49]. Questo vuol dire, in altre parole, che la presenza di una cuspidè non è sufficiente per affermare che al centro dell'ammasso ci sia un IMBH. Ma anche il viceversa non è sufficiente per affermare il contrario. Infatti, la non rilevazione di una cuspidè non esclude a priori la presenza di un IMBH centrale. Come abbiamo già visto, infatti, la sfera di influenza di un IMBH potrebbe rivelarsi decisamente troppo piccola per far sì che la cuspidè si veda.

Oltre ai metodi che si servono della dinamica stellare in ammasso e dei moti propri delle stelle [12], [19], un altro metodo per l'identificazione degli IMBH potrebbe essere quello che si basa sul rilevamento dell'emissione X e radio degli IMBHs in fase di accrescimento [32], [33].

Tali emissioni, però, risultano difficili da osservare nei GCs poichè si tratta di sistemi poveri di gas. Tuttavia, recentemente, le prove dell'esistenza di un IMBH in un GC potrebbero essere fornite dall'osservazione di un evento di disturbo di marea in un ammasso extra-galattico [28]. I dati in X sono stati raccolti dai telescopi in orbita *Chandra* e *XMM-Newton* e successivamente confermati dal *Telescopio Spaziale Hubble*. Attualmente si stanno svolgendo ulteriori indagini osservative per identificare la natura dell'oggetto responsabile dell'evento, ma gli autori dello studio, *Lin D. et al. 2020* [27], sostengono si tratti di un IMBH in procinto di inglobare una stella.

Oltre a questi, esistono anche metodi indiretti per rilevare IMBH al centro di ammassi. Uno tra questi è quello che si basa sull'effetto degli IMBHs sulla segregazione di massa [37]. Nei GCs, ci si aspetta, infatti, che gli IMBHs per la maggior parte del tempo si trovino in sistemi binari con altri oggetti massicci, come buchi neri di origine stellare. In questa

configurazione essi inietterebbero energia nel nucleo del GC, attenuando la segregazione di massa stellare. Anche le binarie primordiali, però, potrebbero essere responsabili dello stesso fenomeno. Ciò porterebbe ad un problema di degenerazione nell'indicatore di segregazione di massa, che potrebbe essere risolto misurando la frazione di binarie nel *core* in maniera indipendente. Questo, però, ha le sue complicazioni dovute ai problemi di risoluzione spaziale di cui si è già parlato. Anche per questo motivo gli studi portati avanti sulla base di tale metodologia non sono ancora stati in grado di rilevare forti candidati di GCs che ospiterebbero un IMBH.

Tutti i metodi presentati riscontrano delle difficoltà nell'identificazione degli IMBHs all'interno degli ammassi. Queste sono legate soprattutto ai fattori che ne influenzano la rilevazione diretta. Gli IMBHs, infatti, risultano così sfuggenti principalmente a causa dei problemi che comporta la piccola dimensione della loro sfera di influenza e, conseguentemente, dei problemi osservativi relativi alla risoluzione spaziale per i *core* degli ammassi. Per tali ragioni risulta interessante approfondire un metodo alternativo e indiretto per la ricerca di IMBHs all'interno di ammassi globulari: quello che si basa sullo studio delle proprietà dinamiche delle Pulsar Millisecondo (*Millisecond Pulsars, MSPs*) [38], [2]. I paragrafi successivi, infatti, saranno dedicati all'approfondimento di tale metodo che è proprio quello su cui si basa questo progetto di tesi.

1.3 Metodo delle MSPs

Le MSPs sono oggetti molto frequenti nei GCs ed essendo caratterizzate da periodi di rotazione estremamente stabili, possono essere utilizzate come strumento per l'identificazione di IMBHs all'interno di ammassi globulari. Le misure dei periodi di MSPs in sistemi binari permettono di ottenere informazioni sulle loro accelerazioni e quindi di ricavarne i

parametri orbitali per mezzo dei quali è possibile capire se esse stiano orbitando attorno ad un IMBH [38].

Inoltre, dalle misure dei *Time of Arrivals* (ToAs) dei segnali emessi dalle MSPs (stavolta non necessariamente in sistemi binari) è possibile calcolare anche le derivate temporali prima e seconda dell'accelerazione, rispettivamente chiamate *jerk* e *snap* [18], [3]. Considerando anche queste quantità è possibile ampliare le informazioni riguardanti le MSPs per identificare la presenza di un IMBH nell'ammasso.

1.3.1 Profili radiali di Jerk e Snap in ammasso

Analiticamente è possibile derivare le relazioni matematiche che descrivono i profili radiali di *jerk* e *snap* di stelle in ammasso [3].

Consideriamo una stella di prova che sperimenta l'attrazione gravitazionale del campo generato da tutte le stelle dell'ammasso, in particolare assumiamo che il GC sia descritto dal profilo di King [23].

L'accelerazione in funzione della distanza r dal centro dell'ammasso, nelle regioni centrali del GC, può essere approssimata come:

$$\mathbf{a}(r) = -4\pi G\rho_c r_c^3 \left[\sinh^{-1}\left(\frac{r}{r_c}\right) - \frac{r}{r_c \sqrt{1 + (r/r_c)^2}} \right] \frac{\mathbf{r}}{r^3} = -|a| \frac{\mathbf{r}}{r} \quad (1.2)$$

dove ρ_c è la densità centrale dell'ammasso ed r_c il raggio del *core*.

Per calcolare il *jerk* è necessario derivare rispetto al tempo l'equazione 1.2, ottenendo:

$$\dot{\mathbf{a}}_K(r) = -\frac{d|a(r)|}{dt} \frac{\mathbf{r}}{r} - |a(r)| \frac{\mathbf{v}}{r} + |a(r)| \frac{(\mathbf{v} \cdot \mathbf{r})\mathbf{r}}{r^3} \quad (1.3)$$

in cui la derivata temporale della norma dell'accelerazione è data da:

$$\frac{d|a(r)|}{dt} = -2\frac{v|a(r)|}{r} + 4\pi Gv\rho_c \left(\frac{1}{1 + (r/r_c)^2} \right)^{\frac{3}{2}} \quad (1.4)$$

con v la norma della velocità.

Nel caso in cui il GC fosse caratterizzato da forti fenomeni di collisioni

stellari, anche i *jerk* potrebbero risultare influenzati dalle stelle vicine. In questo caso, il *jerk* sarebbe caratterizzato dalla seguente distribuzione di probabilità [42]:

$$P(\dot{a}) = \frac{1}{\pi^2} \frac{\dot{a}_0}{(\dot{a}^2 + \dot{a}_0^2)^2} \quad (1.5)$$

in cui $\dot{a}_0$ è il *jerk* caratteristico dato da:

$$\dot{a}_0 = \frac{2\pi\xi}{3} G \langle m \rangle \sigma n \quad (1.6)$$

dove $\xi=3.04$ è una costante numerica, $\langle m \rangle$ è la massa media delle stelle, σ è la dispersione di velocità ed n è la densità numerica delle stelle. La distribuzione dei *jerk* proiettata lungo la linea di vista $\dot{a}_l$ è una distribuzione Lorentziana:

$$P(\dot{a}_l) = \frac{1}{\pi} \frac{\dot{a}_0}{\dot{a}_l^2 + \dot{a}_0^2} \quad (1.7)$$

Se all'interno del GC vi è un IMBH, il *jerk* della stella di prova sarà influenzato dalla massa centrale M ed il profilo sarà:

$$\dot{\mathbf{a}}_M = -GM \left(\frac{\mathbf{v}}{r^3} - 3 \frac{(\mathbf{v} \cdot \mathbf{r}) \mathbf{r}}{r^5} \right) \quad (1.8)$$

dove $\mathbf{r}$ è la distanza dalla massa M e $\mathbf{v}$ è la relativa velocità.

Inoltre l'IMBH crea una sovra-densità stellare caratterizzata da un profilo radiale con *slope* pari a -1.55 [13]:

$$\dot{\mathbf{a}}_{\text{cusp}} = \begin{cases} -\frac{4\pi G}{1.45} r_i^{1.55} \rho_i \left(\frac{\mathbf{v}}{r^{1.55}} - 1.55 \frac{(\mathbf{v} \cdot \mathbf{r}) \mathbf{r}}{r^{3.55}} \right) & \text{for } r < r_i \\ -\frac{4\pi G}{1.45} r_i^3 \rho_i \left(\frac{\mathbf{v}}{r^3} - 3 \frac{(\mathbf{v} \cdot \mathbf{r}) \mathbf{r}}{r^5} \right) & \text{for } r > r_i \end{cases} \quad (1.9)$$

in cui r_i è il raggio di influenza dell'IMBH e ρ_i è il valore di densità a tale raggio.

Derivando rispetto al tempo le equazioni 1.3, 1.8 e 1.9 è possibile ottenere anche i profili radiali degli *snaps* all'interno degli ammassi.

Il contributo del campo gravitazionale dell'intero ammasso è dato dalla:

$$\begin{aligned}\ddot{\mathbf{a}}_K = & -\frac{d^2|a(r)|}{dt^2} \frac{\mathbf{r}}{r} - 2\frac{d|a(r)|}{dt} \frac{\mathbf{v}}{r} + 2\frac{d|a(r)|}{dt} \frac{(\mathbf{v} \cdot \mathbf{r})\mathbf{r}}{r^3} + \\ & + 5|a(r)| \frac{(\mathbf{v} \cdot \mathbf{r})\mathbf{v}}{r^3} - 3|a(r)| \frac{(\mathbf{v} \cdot \mathbf{r})^2\mathbf{r}}{r^5}\end{aligned}\quad (1.10)$$

Analogamente a quanto calcolato per i *jerk*, anche per gli *snap* la presenza di un IMBH al centro dell'ammasso porta a due contributi.

Il contributo che dipende direttamente dalla massa centrale M :

$$\ddot{\mathbf{a}}_M = GM \left(-2a \frac{\mathbf{r}}{r^4} - 6 \frac{(\mathbf{v} \cdot \mathbf{r})\mathbf{v}}{r^5} - 3 \frac{v^2\mathbf{r}}{r^5} + 15 \frac{(\mathbf{v} \cdot \mathbf{r})^2\mathbf{r}}{r^7} \right) \quad (1.11)$$

ed il contributo che determina una sovra-densità. Per quest'ultimo distinguiamo due casi.

Per $r < r_i$:

$$\begin{aligned}\ddot{\mathbf{a}}_{\text{cusp}} = & -\frac{4\pi G}{1.45} r_1^{1.55} \rho_i \left(-0.45 \frac{a\mathbf{r}}{r^{2.55}} - 3.1 \frac{(\mathbf{v} \cdot \mathbf{r})\mathbf{v}}{r^{3.55}} - \right. \\ & \left. -1.55 \frac{v^2\mathbf{r}}{r^{3.55}} + 5.5 \frac{(\mathbf{v} \cdot \mathbf{r})^2\mathbf{r}}{r^{5.55}} \right)\end{aligned}\quad (1.12)$$

e per $r > r_i$:

$$\begin{aligned}\ddot{\mathbf{a}}_{\text{cusp}} = & -\frac{4\pi G}{1.45} r_i^3 \rho_i \left(-2 \frac{a\mathbf{r}}{r^4} - 6 \frac{(\mathbf{v} \cdot \mathbf{r})\mathbf{v}}{r^5} - 5 \frac{v^2\mathbf{r}}{r^5} + \right. \\ & \left. + 15 \frac{(\mathbf{v} \cdot \mathbf{r})^2\mathbf{r}}{r^7} \right)\end{aligned}\quad (1.13)$$

1.3.2 Metodo osservativo

Un ammasso globulare normalmente è caratterizzato in media da un numero di pulsar molto basso, solo in alcuni casi si raggiunge l'ordine di qualche decina di pulsar [43]. Esse si localizzano nelle regioni più interne degli ammassi e di solito circa la metà si trova in sistemi binari.

Dal punto di vista pratico, per identificare IMBHs all'interno di GCs basterebbero le informazioni fornite dalle accelerazioni se avessimo a disposizione un numero consistente di pulsar per ogni ammasso. Ma per ovviare al problema delle poche pulsar normalmente presenti in questi

sistemi, si potrebbero estrarre più informazioni da ognuna di esse considerando anche *jerk* e *snap*. Queste quantità si potrebbero ottenere grazie a periodi prolungati di osservazione delle MSPs nei GCs.

Infatti, le misure delle derivate seconda e terza del periodo di rotazione di un insieme di MSPs in un GC galattico sono fondamentali perché correlano con la componente lungo la linea di vista di *jerk* e *snap*.

I tempi di osservazione richiesti per ottenere tali informazioni e per raggiungere una precisione adeguata sono molto lunghi, dell'ordine di diversi anni. Portando avanti queste campagne osservative in maniera regolare, però, le MSPs potrebbero diventare un buon strumento per l'identificazione di IMBHs al centro degli ammassi.

La derivata prima del periodo è principalmente influenzata dalla componente dell'accelerazione misurata lungo la linea di vista, mentre le derivate seconda e terza del periodo dipendono in maniera diretta da *jerk* e *snap* rispettivamente.

La relazione tra accelerazione lungo la linea di vista a_c e derivata del periodo $\dot{P}$ della MSPs è data da:

$$\left(\frac{\dot{P}}{P}\right)_{\text{meas}} = \left(\frac{\dot{P}}{P}\right)_{\text{int}} + \frac{a_c}{c} + \frac{a_g}{c} + \frac{\mu^2 D}{c} \quad (1.14)$$

in cui $(\dot{P}/P)_{\text{int}}$ è la componente intrinseca dovuta allo *spin-down* della pulsar, a_g è l'accelerazione dovuta al potenziale galattico lungo la linea di vista e l'ultimo addendo rappresenta l'effetto Shklovskii [45], in cui μ è il moto proprio della pulsar, D è la distanza dell'ammasso dal Sole e c è la velocità della luce.

Gli ultimi due termini generalmente possono essere trascurati rispetto al contributo dato dal termine a_c [2]. Purtroppo, però, è molto difficile distinguere gli effetti dell'accelerazione dell'ammasso dallo *spin-down* intrinseco. Qualsiasi lavoro focalizzato sulla misurazione dell'accelerazione in un GC a partire da $\dot{P}$, avrà grandi incertezze dovute all'ignoto *spin-down* intrinseco.

Tuttavia, in maniera indipendente è possibile stimare l’accelerazione delle MSPs in ammasso solo se esse appartengono ad un sistema binario utilizzando l’effetto Doppler.

Per quanto riguarda *jerk* e *snap* la situazione cambia. Le relazioni che legano le derivate seconda e terza del periodo a *jerk*, $\dot{a}_c$, e *snap*, $\ddot{a}_c$, rispettivamente sono date da:

$$\left(\frac{\ddot{P}}{P}\right)_{\text{meas}} = \frac{\dot{a}_c}{c} + \left(\frac{\ddot{P}}{P}\right)_{\text{int}} \quad (1.15)$$

$$\left(\frac{\dddot{P}}{P}\right)_{\text{meas}} = \frac{\ddot{a}_c}{c} + \left(\frac{\dddot{P}}{P}\right)_{\text{int}} \quad (1.16)$$

In questi casi i termini di *spin-down* sono praticamente trascurabili (per ulteriori chiarimenti si rimanda il lettore a [3]). Questo vuol dire che misure di derivata seconda e terza del periodo corrispondono direttamente a misure di *jerk* e *snap*.

1.3.3 Simulazioni e sviluppi con modelli di Machine Learning

Abbate et al. (2019) [3] sviluppano un set di simulazioni a N-corpi di ammassi stellari in cui calcolano *jerk* e *snap* in maniera auto-consistente, trattando le MSPs come particelle di prova.

In primo luogo, dimostrano che la presenza di un IMBH influenza l’andamento di *jerk* e *snap* in funzione della distanza dal centro del GC, specialmente nelle regioni centrali dell’ammasso (Fig. 1.4).

Successivamente, a seguito di un’analisi condotta utilizzando la statistica Bayesiana, gli autori mostrano che con le derivate delle accelerazioni, in particolare con i *jerk*, per identificare un IMBH di massa dell’ordine di $10^2 M_\odot$ è necessario avere a disposizione un numero di MSPs pari a circa 40, di cui 20 devono trovarsi in sistemi binari.

Sulla base di tale lavoro, nasce poi l'idea nuova e sviluppata in questa tesi di prevedere la presenza degli IMBHs all'interno dei GCs tramite un modello di ML. I dati utilizzati nella tesi sono stati ottenuti dalle simulazioni di ammassi globulari condotte dagli autori.

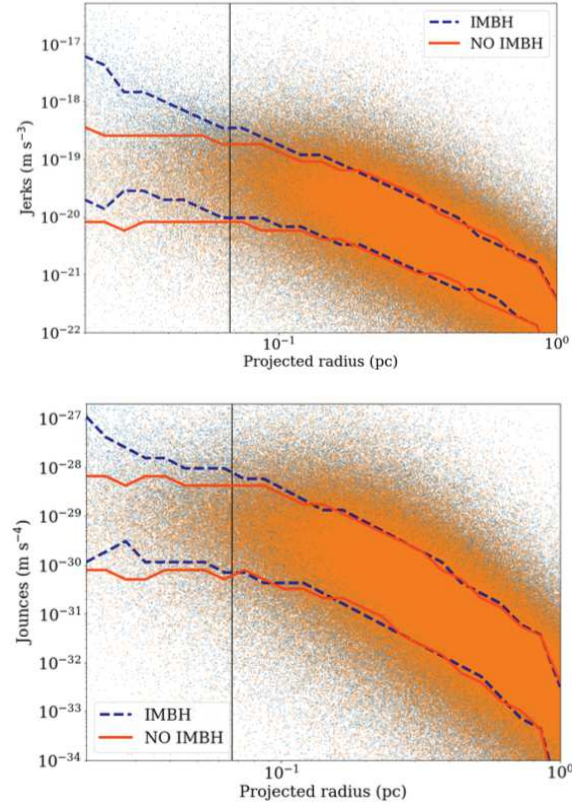


Figura 1.4: Profili radiali di *jerk* (in alto) e *snap* (in basso) ottenuti dalle simulazioni di Abbate et al.2019.

Capitolo 2

Tecniche di Machine Learning interpretabile: alberi decisionali

In questo Capitolo, in seguito ad una rapida introduzione generale al *Machine Learning*, viene presentata una delle tecniche di *Machine Learning interpretabile* e cioè gli alberi decisionali.

Vengono descritti i concetti fondamentali necessari per costruire le strutture ad albero, i problemi più importanti che si possono incontrare durante questa fase ed i metodi che si utilizzano per risolverli.

In particolare, viene descritto l'algoritmo CART [14], ovvero l'algoritmo utilizzato in questo lavoro, ed i metodi su cui esso si basa.

2.1 Introduzione al Machine Learning

Il *Machine Learning*, (*ML*) è una branca dell'Intelligenza Artificiale e si pone l'obiettivo di far apprendere in modo automatico alle macchine

attività svolte da noi esseri umani.

Più precisamente, si dice che un programma impara da una certa esperienza E rispetto ad una classe di compiti T ottenendo una *performance* P , se la sua *performance* P nel realizzare i compiti T , migliora con l'esperienza E [31].

In altre parole, se un programma migliora lo svolgere di un *task* rispetto a un'esperienza passata, si dice che ha imparato.

Questo avviene tramite l'apprendimento automatico. Esso può essere suddiviso in due importanti categorie: apprendimento supervisionato e apprendimento non supervisionato.

Gli algoritmi di apprendimento automatico supervisionato sono sequenze di operazioni che utilizzano i dati di allenamento (*training set*) ricevuti in input per produrre un modello capace di risolvere un problema di classificazione o di regressione su dati di test (*test set*) mai visti in precedenza, con una *performance* che aumenta in funzione della quantità di dati di allenamento ricevuti. In primo luogo, quindi, l'algoritmo lavora su un sottoinsieme del *dataset* chiamato *training set*. Una volta costruito il modello, questo poi viene utilizzato per riconoscere e analizzare dati mai visti (chiamati dati del *test set*).

Durante l'apprendimento supervisionato si hanno a disposizione sia i dati di input (X , matrice delle *features* composta dalle variabili indipendenti che usiamo per la predizione), sia i dati di output (Y , composto dalle variabili dipendente che vogliamo predire chiamate *labels*). In questo caso si utilizza un algoritmo che apprende la funzione f che dall'input genera l'output: $Y = f(X)$.

L'obiettivo è approssimare la funzione in modo che quando si ha un nuovo dato di input l'algoritmo sia in grado di prevedere il valore di output generato per quel dato.

I problemi di apprendimento supervisionato possono essere distinti in:

- Classificazione: si tratta di un problema discreto, cioè la *label* è

una variabile categorica (sì/no, vero/falso, 0/1/2...). Ad esempio, in campo medico, si potrebbe voler determinare in base ai risultati quantitativi di una biopsia se un caso di tumore è benigno o maligno;

- Regressione: si tratta di un problema continuo, cioè la *label* è una variabile numerica. Ad esempio, il prezzo più probabile di una casa che si vuole prevedere sulla base dei metri quadri e della zona di interesse.

Nell'apprendimento non supervisionato, invece, si ha solo la variabile di input X e nessuna variabile di output corrispondente. Pertanto, gli algoritmi di apprendimento non supervisionato cercano di trovare una struttura nel *dataset*.

In questo caso i problemi di apprendimento possono essere suddivisi in:

- Raggruppamento: anche detto *clustering*, viene utilizzato quando è necessario raggruppare i dati che presentano caratteristiche simili. Per esempio un assicuratore potrebbe voler individuare clienti che corrispondono a profili di rischio simili tra loro, sulla base di caratteristiche quali: l'età, il genere, indicatori dello stato di salute, ecc...;
- Associazione: è un problema dove si vogliono scoprire regole che descrivono grandi porzioni di dati; si ha come obiettivo quello di trovare schemi frequenti, associazioni, correlazioni o strutture casuali tra un insieme di oggetti in un *dataset* relazionale. Un'applicazione è quella del *market basket analysis*. Si tratta di analisi di transazioni commerciali che possono produrre informazioni per determinare regole ricorrenti che pongono in relazione l'acquisto di uno o più prodotti con altri. Ad esempio, se un cliente compra il latte qual è la probabilità che compri anche i cereali? Sulla base di queste informazioni si possono progettare azioni promozionali o posizionare gli articoli sugli scaffali;

- Riduzione della dimensionalità: è un problema in cui si vuole individuare un numero ridotto di *features* rappresentative delle caratteristiche dei dati all'interno di un grande campione di *features* inizialmente disponibile. Ad esempio, uno psicologo potrebbe utilizzare dati di un campione di studenti delle scuole superiori per costruire un indicatore di abilità linguistica e uno di abilità numerica combinando i voti ricevuti nelle varie materie. Questo consentirebbe la visualizzazione di una pagella come un punto nel piano definito da queste due variabili, invece di trovarsi ad affrontare il problema più complesso di visualizzare contemporaneamente i voti di tutte le materie.

In questo lavoro di tesi ho utilizzato un metodo supervisionato di classificazione: gli alberi decisionali. L'importanza di questo metodo sta nel fatto che gli alberi possono essere interpretati, in quanto è possibile rappresentarli graficamente. Nelle prossime Sezioni, infatti, saranno approfonditi gli aspetti teorici che li riguardano.

2.2 Concetti preliminari: notazione e struttura del dataset

In questa Sezione vengono esposti alcuni concetti preliminari che riguardano gli algoritmi di classificazione, necessari per entrare nel merito dei metodi utilizzati per questo lavoro.

Il *dataset* sul quale avviene la fase di addestramento viene chiamato *training set* e il modo più semplice per descriverlo è mediante una matrice $\mathbf{X} \in \mathbb{R}^{n \times m}$. Le righe della matrice sono i *records*, ovvero gli esempi o le osservazioni e le colonne sono le *features*, cioè le caratteristiche multiple aventi per ogni esempio.

Per un algoritmo di classificazione, essendo un metodo di addestramento supervisionato, il *dataset* è caratterizzato anche dal vettore Y delle

labels. Ogni *label*, ovvero ogni etichetta, corrisponde ad una classe e sono le variabili target che si vogliono predire (tab. 2.1).

	feat. 1 (A_1)	feat. 2 (A_2)		feat. m (A_m)	classe
record 1					
record 2					
record 3					
.....					
.....					
record n					

Tabella 2.1: Schema della struttura di un generico *dataset* di apprendimento nel caso in cui si stiano usando algoritmi di classificazione.

La classificazione ha come scopo quello di analizzare i dati di input sviluppando un modello in grado di predire la classe di appartenenza dei *records* in base alle *features* presenti nei dati. In genere, partendo dall'utilizzo di insiemi esistenti e già classificati, l'algoritmo cerca di identificare alcune regolarità che caratterizzano le varie classi.

2.3 Alberi decisionali: concetti generali

Gli alberi decisionali sono strutture molto conosciute nell'ambito degli algoritmi supervisionati, in quanto permettono di classificare in modo semplice degli oggetti in un numero finito di classi.

L'utilizzo degli alberi decisionali offre numerosi vantaggi:

- sono di facile interpretazione, specie se non sono molto profondi;
- ottengono una buona accuratezza su gran parte dei problemi reali di classificazione su dati tabulari;
- sono robusti rispetto al rumore e alla ridondanza tra le *features*;

- possono essere costruiti efficientemente ed essere visualizzati.

Gli alberi vengono costruiti suddividendo i *records* in sottoinsiemi in base alle relazioni che legano le variabili target, che si cercano di prevedere, alle *features* utilizzate come predittori. In particolare, questo permette di costruire un modello rappresentato da un insieme di regole ottenute ponendo una serie di domande (test). Queste sono mirate sui valori delle *features* e tipicamente consistenti in un confronto tra il valore di una data *feature* e una soglia numerica. Ogni volta che si riceve una risposta viene posta la domanda successiva in modo che sia attinente al risultato ottenuto. Il processo viene iterato fino all'ottenimento della classe di ciascun *record*. La serie di domande, e le relative risposte sono organizzate in una struttura ad albero.

Fondamentalmente si tratta di una struttura semplice composta da nodi, rami e foglie (anche dette nodi terminali) che si sviluppa a partire dal nodo radice. Ogni nodo corrisponde ad una decisione basata sul confronto di una *feature* con una costante. Essi sono collegati dai rami che identificano i livelli di parentela tra i diversi nodi (il nodo genitore rappresenta una decisione presa a monte rispetto al nodo figlio) e che forniscono gli strumenti per la costruzione di regole necessarie per classificare un oggetto. Infine, i risultati sono identificati dalle foglie. Dato un albero decisionale si possono ottenere predizioni rispetto a nuovi dati semplicemente percorrendo l'albero dalla radice verso le foglie, seguendo di volta in volta il ramo corrispondente al risultato del confronto effettuato in ciascun nodo.

Per fare maggiore chiarezza su come si costruisce un albero decisionale e sulla sua struttura, di seguito viene riportato un semplice esempio.

Supponiamo che una compagnia di assicurazioni voglia identificare il legame che lega le classi di rischio, in cui vengono suddivisi i clienti, con la loro età anagrafica e il tipo di vettura posseduto. Lo studio si deve basare su un gruppo di clienti già classificati nella corrispettiva classe.

Il *training set* T a disposizione è rappresentato in tabella 2.2 e l'insieme

delle classi è $\Gamma = \{A, B\}$, in cui A identifica un Alto rischio e B un Basso rischio.

Rid	Età	Tipo di Auto	Rischio
1	23	Berlina	A
2	18	Sportiva	A
3	43	Sportiva	A
4	68	Berlina	B
5	32	Furgone	B
6	20	Berlina	A

Tabella 2.2: *Dataset* d'esempio per classificare clienti di una compagnia di assicurazioni in opportune classi di rischio.

In figura 2.1 viene riportato un possibile albero decisionale per l'esempio in questione. Si noti come i nodi interni corrispondano a test sulle *features* e come i nodi foglia vengano etichettati con la classe di maggioranza per la rispettiva foglia.

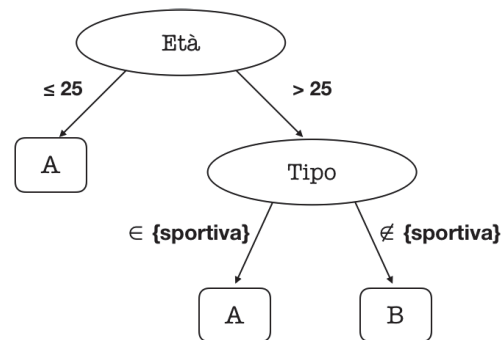


Figura 2.1: Esempio di albero decisionale per una compagnia di assicurazioni.

I test sulle *features*, effettuati nei nodi interni, differiscono a seconda del tipo di dati. La suddivisione del *dataset*, dovuta ad esiti diversi nei test, è definita *split*.

In questo esempio gli *split* sono binari, in quanto per ogni nodo si hanno a disposizione due possibilità.

2.4 Algoritmo CART

Esistono in letteratura diversi algoritmi di classificazione che fanno uso degli alberi decisionali. Quello che è stato utilizzato in questo lavoro è l'algoritmo CART.

L'algoritmo CART, *Classification and Regression Trees* [14], è uno degli algoritmi più conosciuti ed utilizzato per lo sviluppo di alberi decisionali e può lavorare sia come un classificatore che come un regressore.

Uno dei suoi punti di forza è la sua semplicità: esso opera mediante degli *split* binari (ad ogni nodo corrispondono due soli rami) su una singola variabile in modo ricorsivo, quindi, classificare un campione può richiedere solo pochi semplici passi.

Nonostante la sua semplicità, è comunque in grado di ottenere risultati migliori di diversi altri metodi, su *dataset* complessi, non lineari e composti da molte variabili.

Questo è senz'altro un valore aggiunto al fatto che gli alberi decisionali siano di così facile ed intuitiva interpretazione, per questo motivo vengono utilizzati negli ambiti più svariati. Quelle appena discusse sono anche le motivazioni primarie che hanno spinto all'utilizzo di tale metodologia per lo sviluppo di questo progetto di tesi. Trattandosi di un *dataset* non ancora mai sottoposto ad uno studio con metodi di ML, gli alberi decisionali con l'algoritmo CART hanno fornito una prima analisi esplorativa dei dati fornendo già dei buoni risultati (Cap. 4).

Il metodo CART è molto efficiente anche perché non richiede in ambito applicativo di sperimentare trasformazioni delle variabili indipendenti (logaritmi, radici quadrate, elevamento a potenza, ecc). Tali trasformazioni, in quanto monotone, non modificano i risultati a meno che lo *split*

sia basato su combinazioni lineari delle variabili. Questo in genere non succede in un albero convenzionale, dove le *features* vengono valutate indipendentemente in ciascun nodo. Inoltre, esso può utilizzare la stessa variabile in punti differenti dell'albero.

Il CART è stato creato con lo scopo di trattare dati dotati di una struttura complessa. E' estremamente robusto all'effetto degli *outliers* e può utilizzare congiuntamente variabili di tipo categorico e continue.

L'algoritmo CART produce alberi ottimali, utilizzando strumenti sofisticati per stabilirne l'accuratezza (Sez.2.11). E, infine, oltre ad essere alla base di altri algoritmi che generano alberi più complessi, esso può essere utilizzato a supporto di altri tipi di modelli.

L'algoritmo si sviluppa in due fasi:

- Fase di generazione dell'albero;
- Fase di riduzione della complessità.

Vedremo ora quali sono in generale i criteri per la costruzione di alberi decisionali e su quali in particolare il CART si basa.

2.5 Criteri per la costruzione di alberi decisionali

Per costruire una struttura ad albero efficace occorre seguire tre *step*:

- Selezionare una regola per lo *split* per ogni nodo, ciò significa determinare le variabili indipendenti ed i rispettivi valori soglia, che saranno usati per partizionare il *training set*;
- Determinare quali nodi sono terminali, quindi decidere quando "fermarsi";
- Assegnare una classe ad ogni nodo terminale.

2.6 Regole di Splitting

Definire le regole di *splitting* vuol dire identificare, per ogni nodo, una variabile sulla quale effettuare il test sulla base del rispettivo valore soglia utilizzato per il partizionamento del *training set*.

Il punto è quello di scegliere in ogni nodo la migliore variabile di *split*, che garantisca la suddivisione dei *records* presenti nel nodo in sottogruppi il più possibili omogenei al loro interno ed eterogenei tra loro in termini dei valori delle rispettive *labels*.

Occorre, quindi, effettuare una sorta di valutazione della bontà degli *split*. A tal proposito gli autori dell'algoritmo CART hanno sviluppato un *framework* metodologico, introducendo il concetto generico di *impurità*.

Intuitivamente, quando dividiamo i dati, vorremmo che la regione corrispondente a ciascun nodo foglia sia "pura", ovvero che la maggior parte dei dati in questa regione provenga dalla stessa classe e quindi che ci sia una classe dominante.

Consideriamo il seguente esempio [8] mostrato in figura 2.2. Qui ab-

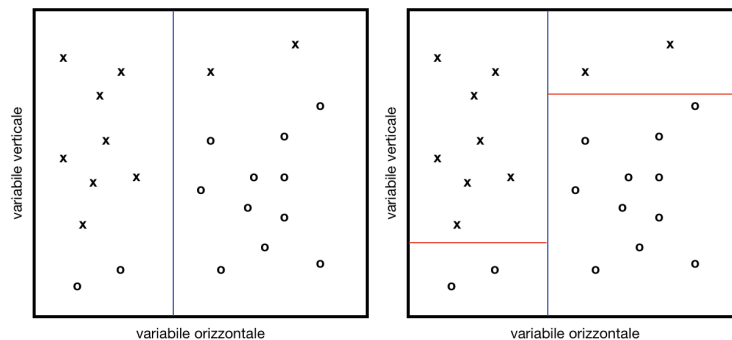


Figura 2.2: Un semplice esempio di *splitting*. La figura a sinistra mostra il primo *split* (linea blu), mentre quella a destra mostra anche il secondo (linea rossa).

biamo due classi: la classe **x** e la classe **o**. Le variabili di input sono la variabile orizzontale e quella verticale. Il primo *split* viene fatto controllando se la variabile orizzontale è al di sopra o al di sotto di una soglia (la divisione è indicata dalla linea blu).

Potremmo considerare questo *split* una buona divisione perché il lato sinistro è quasi puro in quanto la maggior parte dei punti appartiene alla classe **x** e solo due punti appartengono alla classe **o**. Il viceversa vale per il lato destro.

Proseguendo ad un livello più in basso nell'albero, vediamo che sono state create altre due divisioni (linee rosse). La regione in alto a sinistra (o il nodo foglia) contiene solo la classe **x**, così come quella in alto a destra. Invece le regioni in basso a sinistra ed in basso a destra contengono solo la classe **o**.

A questo punto non è necessario effettuare altri *split* perché tutte le foglie sono pure al 100%.

2.6.1 La funzione di impurità

La misura di impurità può essere ricavata a partire dalla cosiddetta funzione di impurità ϕ .

Essa misura l'entità della purezza di una regione contenente osservazioni appartenenti a classi possibilmente diverse.

Supponiamo che il numero di classi sia k . Quindi la funzione di impurità è definita sul set di k -tuple $p_1, p_2, \dots, p_k$, che sono le probabilità di ogni osservazione di appartenere ad una certa classe. Vale che $p_j \in [0,1] \forall j = 1, \dots, k$ e $\sum_j p_j = 1$.

Durante l'addestramento non si conoscono le reali probabilità; si utilizzano, infatti, le percentuali di osservazioni in classe 1, classe 2, classe 3 e così via, in base al set di dati di addestramento.

La funzione di impurità può essere pensata come una funzione avente

valori tra 0 e 1 che fornisce una misura di quanto le osservazioni siano correttamente distribuite nelle classi.

Si definisce, per un generico nodo t , la misura di *impurità* $i(t)$ come segue:

$$i(t) = \phi(p(1 | t), p(2 | t), \dots, p(k | t)) \quad (2.1)$$

ove $p(j | t)$, con $j = 1, \dots, k$, è la probabilità stimata a posteriori di ottenere la classe j per un'osservazione nel nodo t .

L'impurità di un nodo è massima quando tutte le classi sono presenti nella stessa proporzione, mentre è minima quando il nodo contiene casi appartenenti ad un'unica classe. La misura di impurità viene usata per decidere quale *split* fare in un dato nodo, determinando in sostanza in che modo si fa crescere l'albero.

Misure di impurità più comuni sono:

- L'indice di eterogeneità di Gini: $\text{Gini}(t) = 1 - \sum_j p_j^2(t)$
- L'entropia: $H(t) = - \sum_j p_j \log_2 p_j(t)$
- Il tasso di errata classificazione: $r(t) = 1 - \max \{p_j(t) : 1 \leq j \leq k\}$

Generalmente, in particolare negli algoritmi che lavorano con *splitting* binari come CART, viene utilizzato l'indice di Gini.

Tale misura può essere interpretata come la stima della probabilità che un'osservazione scelta casualmente nel nodo t sia assegnata alla classe errata.

Essa tende a dare luogo a suddivisioni bilanciate dal punto di vista del numero di casi inviati dallo *split* nei due nodi figli, ovvero tende ad evitare i cosiddetti *small splits*.

2.7 Dichiarazione dei nodi terminali

A questo punto, per completare la costruzione dell'albero è necessario passare per gli ultimi due *step*.

Bisogna capire, in primo luogo, quando arrestare la "crescita" dell'albero e quindi identificare i nodi terminali. Infine è necessario assegnare una classe ad ogni nodo terminale.

2.7.1 Criteri di arresto

Nella maggior parte dei casi, se si lasciasse sviluppare un albero fino alle foglie estreme, si avrebbe una struttura molto grande e caratterizzata dal concatenarsi di numerose condizioni. Così facendo, il vantaggio della semplicità interpretativa, caratteristica degli alberi decisionali, verrebbe perso e si andrebbe incontro al problema dell'*overfitting* (approfondito nella Sezione 2.8).

La dimensione eccessiva sarebbe dovuta al numero di nodi terminali o, equivalentemente, al numero di suddivisioni che essi rappresentano. Per questa ragione, per ovviare al problema, si può decidere di interrompere l'espansione dell'albero sulla base di determinati criteri.

Alcuni criteri comuni sono:

- continuare la crescita finché si raggiunge una certa profondità predefinita;
- continuare finché il numero di osservazioni in ciascun nodo terminale non superi una certa soglia;
- continuare finché tutti i nodi terminali non sono puri, cioè contengono solo una classe; ciò equivale a far crescere interamente l'albero.

Queste soglie vanno decise a priori e influenzano direttamente la dimensione dell'albero. Per questo motivo, sebbene sia questo il metodo più semplice, risulta piuttosto inefficiente perché rischia di sfoltire l'albero troppo o troppo poco rispetto al necessario, in base a quanto sia restrittivo il criterio d'arresto.

2.7.2 Assegnazione della classe al nodo terminale

Nella fase di assegnazione di una delle classi ai nodi terminali si possono presentare tre situazioni diverse:

- il nodo contiene solo osservazioni appartenenti alla stessa classe;
- il nodo contiene osservazioni appartenenti a classi diverse, ma una di queste presenta una proporzione di osservazioni maggiore delle altre;
- Il nodo contiene osservazioni appartenenti a classi diverse e nella stessa proporzione.

Nel primo caso l'assegnazione viene effettuata precisamente senza alcun tipo di indecisione e nel secondo caso il nodo viene assegnato alla classe che presenta più osservazioni. Nell'ultimo caso, invece, si presenta una situazione di massima incertezza, in quanto la probabilità delle classi risulta identica per ciascuna di esse. Pertanto si ricorre ad un tipo di assegnazione casuale salvo l'intervento del ricercatore che può effettuare un'assegnazione diversa sulla base delle sue conoscenze.

2.8 Problema dell'overfitting

Alla fase di costruzione dell'albero decisionale, per far sì che esso si comporti come un buon classificatore, segue la fase di riduzione della complessità dell'albero. Ma prima di approfondire la descrizione che riguarda questa seconda fase, bisogna introdurre un problema con il quale, quasi sempre, tutti gli algoritmi di ML si confrontano: il problema dell'*overfitting*.

Il metodo di classificazione è un processo che si sviluppa in due fasi: la fase di apprendimento e la fase di predizione.

Nella fase di apprendimento, il modello viene sviluppato sulla base di

dati di allenamento conosciuti, che abbiamo chiamato *training set*.

Successivamente, la fase di predizione viene eseguita su un insieme di dati totalmente nuovi per il modello, il cosiddetto *test set*.

In pratica, quindi, il modello apprende informazioni, relazioni e collegamenti durante la prima fase e poi applica tutto ciò che ha acquisito su un nuovo insieme di dati per fornire una predizione.

Quello che normalmente accade nei metodi di classificazione, come ad esempio gli alberi decisionali, è che nelle prime prove di addestramento il modello porta a delle prestazioni molto buone sul *trainig set* e alquanto scadenti sul *test set*. Vuol dire che il modello impara caratteristiche del *training set* che sono proprie solo di quest'ultimo e non si applicano altrove. Ciò non permette, quindi, di generalizzare quanto appreso su un nuovo gruppo di dati. Per esempio, in un certo *training set*, per puro caso, tutti i fumatori sono pisani. Nella realtà, anche persone di altre città fumano con la stessa probabilità e quindi la città di provenienza è pressoché inutile per prevedere se una persona fuma o meno. Ma, specialmente in *training set* piccoli, correlazioni spurie di questo tipo possono verificarsi abbastanza di frequente.

Per "generalizzazione", quindi, si intende l'abilità di una macchina di portare a termine in maniera accurata esempi o compiti nuovi, che non ha mai affrontato, dopo aver fatto esperienza su un insieme di dati di apprendimento [9].

Questo problema è noto con il termine *overfitting*, ovvero si tratta di un problema di sovradattamento da parte del modello ai dati *trainig set*. Si verifica quando il modello viene addestrato eccessivamente su un set di dati rumorosi.

Pertanto, a causa di questo problema, quando si sviluppa un albero decisionale, è importante capire quando fermarsi. Se un albero crescesse troppo, fino alle foglie più estreme, oltre a perdere la sua importante caratteristica di essere interpretabile, si andrebbe incontro al problema dell'*overfitting*. Questo è evidente se, ad esempio, avessimo un albero

con una sola osservazione per foglia, raggiungendo quindi una purezza perfetta in ciascuna foglia sul *training set*. Se si estendesse al massimo un albero, infatti, tale obiettivo si potrebbe sempre raggiungere per qualunque scelta dei dati di *training*, ma negli ultimi *split* l'albero non avrebbe appreso alcuna informazione utile dai dati, adattandosi solo alle idiosincrasie del *training set*. Pertanto, risulterebbe probabile che la *performance* predittiva su dati non visti sia molto inferiore, nonostante la predizione nominalmente perfetta in *training*.

Una tecnica per risolvere tale problema per gli alberi decisionali è la tecnica del *pruning*. In breve, essa consiste nell'ottenere da un albero il più piccolo "sottoalbero" che, di fatto, non comprometta l'accuratezza della classificazione. In particolare, nell'algoritmo CART, la tecnica utilizzata è quella del *Cost-Complexity Pruning* che verrà approfondita nella prossima Sezione.

2.9 Cost-Complexity Pruning

Per ridurre la complessità dell'albero affinché possa essere effettivamente utile nel classificare *records* e nel fornire regole efficaci, è necessario sfoltire la ridondanza dell'albero, ovvero eliminare i rami meno significativi.

La tecnica che l'algoritmo CART utilizza prende il nome di *Cost-Complexity Pruning*. Consiste nel far sviluppare l'intero albero e poi estrarre da esso il più piccolo "sottoalbero". Quest'ultimo porterà a delle migliori previsioni perché vi sono stati "tagliati" i rami corrispondenti a *splitting* eccessivamente specifici sul *training set* e, quindi, superflui. Infatti *pruning* sta proprio per "potatura".

Questo processo, appunto, non peggiora l'accuratezza delle previsioni, ma bensì, in genere, le migliora, in quanto, poi, il modello sarà capace di classificare efficientemente anche dati totalmente nuovi.

2.9.1 Notazione e metodo

Ora viene introdotta la notazione necessaria per la trattazione della tecnica del *Cost-Complexity Pruning*.

Si definiscono:

- T_{max} dimensione di massima crescita dell'albero T ;
- t i nodi genitori;
- t' i nodi figli, cioè un nodo collegato tramite un percorso al nodo t ;
- T_t un ramo dell'albero T con nodo radice t .

Un metodo di potatura efficiente dovrebbe garantire che la ricerca della sottostruttura ottimale possa essere trattabile computazionalmente.

L'algoritmo del *Cost-Complexity Pruning* è parametrizzato dal parametro $\alpha \in \mathbb{R}$, $\alpha \in [0,1]$, noto come *parametro di complessità*. Questo parametro viene utilizzato per definire la misura *cost-complexity* $R_\alpha(T)$ per ogni sottostruttura $T < T_{max}$ come:

$$R_\alpha(T) = R(T) + \alpha|\tilde{T}| \quad (2.2)$$

dove $R(T)$ è il tasso di errata classificazione totale dei nodi terminali (definito in 2.6.1) e $|\tilde{T}|$ è il numero di nodi terminali (o nodi foglia) che definisce la complessità dell'albero T . Infatti, maggiore è il numero di nodi foglia che l'albero contiene, maggiore è la complessità dell'albero perché vi è maggiore flessibilità nel partizionare lo spazio in parti più piccole e quindi maggiori possibilità di adattare i dati di addestramento. Durante la potatura dell'albero bisogna minimizzare la funzione $R_\alpha(T)$ e la sottostruttura che alla fine verrà selezionata dipende dal parametro α . Infatti, se $\alpha = 0$ verrà scelto l'albero più grande perché il termine di complessità viene essenzialmente eliminato. Per α tendente all'infinito, invece, verrà selezionato l'albero di dimensione 1, cioè un singolo nodo

radice.

Si può dimostrare [8] che per ogni α esiste ed è unica una sottostruttura che minimizza $R_\alpha(T)$. Inoltre, si può dimostrare anche che le sottostrutture che minimizzano $R_\alpha(T)$ al crescere di α sono annidate. Ovvero, il sottoalbero dell' α successivo è compreso in quello precedente e, quindi, vale che $T_1 > T_2 > \dots > t_1$, dove i pedici rappresentano il numero progressivo degli α . Praticamente, a partire dalle foglie l'algoritmo procede spostandosi verso la radice valutando in ogni nodo la *cost-complexity* al variare di α .

Per fare maggiore chiarezza sul metodo di taglio, in primo luogo, è necessario estendere la definizione 2.2 a un nodo e ad un singolo ramo fuoriuscente dal nodo:

- per qualsiasi nodo t appartenente ad una sottostruttura la 2.2 diventa:

$$R_\alpha(t) = R(t) + \alpha$$

in quanto ci si sta riferendo proprio ad un nodo, in questo caso il termine di complessità $|\tilde{T}|$ è pari a 1;

- per qualsiasi ramo T_t , invece la 2.2 diventa:

$$R_\alpha(T_t) = R(T_t) + \alpha |\tilde{T}_t|$$

in cui il termine $R(T_t)$ è calcolato considerando tutti i nodi terminali che discendono dal nodo t attraverso il ramo T_t (Fig. 2.3).

In ogni nodo, facendo variare α tra 0 e 1 in maniera crescente, l'algoritmo confronta il *cost-complexity* calcolato nel nodo t con quello calcolato per il ramo T_t . Affinchè l'algoritmo decida di non tagliare un ramo, lo *split* che avviene per mezzo di esso, deve essere efficiente. Questo si verifica quando lo *split* riduce l'errore sulla classificazione. Pertanto, finchè $R_\alpha(T_t) < R_\alpha(t)$ il ramo viene tenuto, ma non appena i due valori si eguagliano esso

viene tagliato. Per cui, la condizione di taglio è data da $R_\alpha(T_t) = R_\alpha(t)$, da cui:

$$\alpha = \frac{R(t) - R(T_t)}{|\tilde{T}_t| - 1} \quad (2.3)$$

L'algoritmo lavora percorrendo a ritroso l'albero, dalle foglie verso la radice.

Inizialmente, quindi, considera tutti i nodi che precedono le foglie e, facendo variare α , il primo nodo trovato che soddisfa la condizione 2.3 viene tagliato. A questo punto, l'algoritmo salva in memoria l' α appena ricavato e poi procede al passo successivo considerando l'albero appena ottenuto come albero iniziale. L'algoritmo procede in questo modo sull'intero albero.

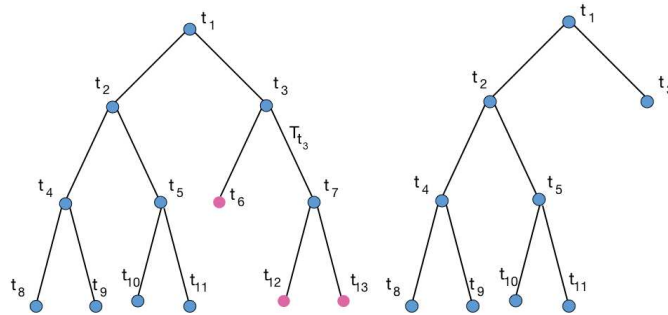


Figura 2.3: Un semplice schema di una struttura ad albero. I nodi terminali del ramo T_{t_3} sono t_6 , t_{12} e t_{13} (a sinistra). In t_3 si verifica la condizione 2.3, per cui l'intero ramo T_{t_3} viene tagliato (a destra).

Pertanto, la scelta del parametro α nella fase di potatura dell'albero può essere presa direttamente dall'utente in base a conoscenze specifiche legate al problema, oppure sperimentalmente, tramite una fase di validazione, approfondita nella Sezione successiva.

2.10 Training, validation e test sets

Per identificare i parametri da fornire ai modelli di ML è di uso comune suddividere il *dataset* in tre sottogruppi: *training set*, *validation set* e *test set*.

In questo caso particolare, per effettuare un *pruning* efficiente, dal punto di vista pratico, è necessario identificare il parametro di complessità α al quale corrisponderebbe la miglior sottostruttura dell'albero T . Infatti, come spiegato nel paragrafo precedente, ad ogni passo l'algoritmo *Cost-Complexity Pruning* salva in memoria il parametro α che soddisfa la relazione 2.3 con il relativo modello ad albero. Sarà l'utente, in una fase successiva, a scegliere tra questi la miglior sottostruttura, cioè quella per cui si ottiene la miglior accuratezza nella classificazione. Per questo motivo si divide il *dataset* in tre sottogruppi.

Il *training set* è l'insieme di dati con i *records* e le relative *labels* e serve al classificatore durante la fase di allenamento per apprendere gli schemi che possono essere utilizzati successivamente per prevedere le *labels* dei nuovi dati. In questa fase è importante che non venga effettuata nessuna scelta sul modello.

Il *validation set*, invece, è utile per determinare con quali parametri l'algoritmo di classificazione avrebbe le migliori prestazioni. Nel caso degli alberi decisionali la fase di validazione è necessaria per scegliere il miglior parametro di complessità α .

Una volta ottenuto il miglior classificatore bisogna valutare le sue prestazioni su un set di dati completamente nuovo per lui: il *test set*.

La scelta su come effettuare la divisione del *dataset* è individuale, ma generalmente il *test set* rappresenta circa il 20 – 30% dei dati. La restante parte del *dataset* viene suddiviso nel seguente modo: circa il 70 – 80% dei dati rimasti costituirà il *training set* e circa il 20 – 30% farà parte del *validation set*.

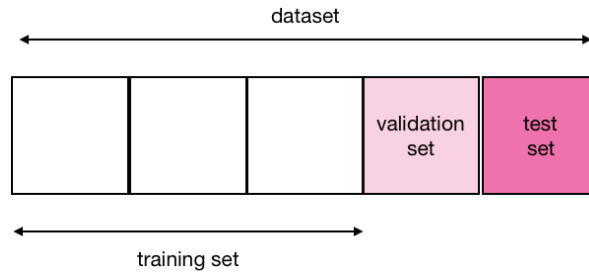


Figura 2.4: Schema per visualizzare la divisione dei dati nei tre sottogruppi: *training set*, *validation set* e *test set*.

2.11 Qualità delle previsioni

Per qualsiasi algoritmo di ML, dopo aver ottenuto le previsioni, bisogna valutarne la qualità. Per quantificare la bontà delle *performance* possono essere utilizzate diverse metriche.

In questo lavoro, per valutare le prestazioni della classificazione da parte degli alberi decisionali è stata utilizzata l'*accuratezza* e altre metriche come *precisione*, *richiamo* ed *F-score* (o *F-1*), che misurano aspetti specifici delle prestazioni di classificazione. Esse verranno approfondite successivamente.

2.11.1 Accuratezza degli alberi decisionali

Per definire l'accuratezza di un albero decisionale, in primo luogo, introduciamo una notazione. Indichiamo con:

- $\bar{n}_t$ il numero totale dei *records* del *test set* che finiscono nel nodo terminale t ;
- n_t il numero di *records* classificati correttamente in t .

Ad ogni nodo terminale t è possibile associare un'accuratezza a_t data dalla frazione dei *records* classificati correttamente nella data foglia t :

$$a_t = \frac{n_t}{\bar{n}_t} \quad (2.4)$$

L'accuratezza nella classificazione associato all'intero albero è data dalla somma delle accuratze calcolate su tutti i t , pesate rispetto alla probabilità che un *record* possa finire su ciascuna foglia.

Se il numero totale di *records* è N , allora, la probabilità che un certo *record* finisca nella foglia t è data da $\frac{\bar{n}_t}{N}$. Pertanto l'accuratezza A dell'albero risulta:

$$A = \sum_t \frac{\bar{n}_t}{N} a_t \quad (2.5)$$

La situazione ideale si ha quando $A = 1$, in cui il modello è in grado di classificare perfettamente tutti i *records*.

Questa metrica di solito lavora bene su un set di dati bilanciato, ovvero quando si ha un numero simile di campioni appartenenti a ogni classe. L'accuratezza, però, non è il miglior metodo per valutare le prestazioni di un classificatore.

2.11.2 Matrice di confusione: precisione, richiamo e F-score

Uno strumento intuitivo per valutare le *performance* di un modello di classificazione è la *matrice di confusione*. Per analizzare le informazioni che essa fornisce bisogna definire le metriche che utilizza e cioè la *precisione*, il *richiamo* e la *F-score* (o *F-1*).

Supponiamo di avere due classi: la classe P dei positivi e la classe N dei negativi.

Allora, le previsioni ottenute da un modello su un *test set* possono essere definite nei seguenti modi:

- **Veri Positivi** (TP), sono i dati test realmente etichettati come P e che il modello ha previsto correttamente come P ;
- **Falsi Positivi** (FP), sono i dati test etichettati come N , ma per i quali il modello ha predetto essere P ;
- **Falsi Negativi** (FN), sono i dati test appartenenti alla classe P , ma per cui il modello ha predetto come N ;
- **Veri Negativi** (TN), sono i dati test della classe N per i quali il modello ha predetto correttamente collocandoli nella classe N .

A partire da queste definizioni è possibile introdurre le metriche:

- **Precisione**, $P = \frac{TP}{TP+FP}$, è il numero di TP diviso per il numero di classificazioni positive, cioè TP e FP, ad esempio, il numero di persone veramente malate su quelle che un dato test ritiene malate;
- **Richiamo**, $R = \frac{TP}{TP+FN}$, è il numero di TP diviso per il numero di tutti i *records* positivi, cioè TP e FN, ad esempio, il numero di persone malate individuate da un test sul totale delle persone veramente malate;
- **F-score**, $F_1 = 2 \cdot \frac{P \cdot R}{P+R}$, è una media armonica tra precisione e richiamo. La media armonica è minore della media aritmetica, pertanto, per ottenere una buona *F-score* bisogna avere sia un buon valore di richiamo che di precisione.

Un valore vicino a 1 della *F-score* segnala la capacità del modello di classificare correttamente la maggior parte dei dati, ottenendo sia una buona precisione che un buon richiamo (Fig. 2.5). D'altra parte ottenere un buon richiamo con una pessima precisione è semplice: basta dichiarare tutti malati, riducendo così a zero il numero di falsi negativi al costo di avere molti falsi positivi. Ma il classificatore che ne risulta è inutile. Per questo motivo, il richiamo da solo non è una buona metrica. Per ragioni

analoghe (si potrebbe dichiarare che nessuno è malato, portando i falsi positivi a zero) la precisione da sola è a sua volta una metrica non buona. Combinandole, la *F-score* garantisce una misura più bilanciata tra questi due aspetti.

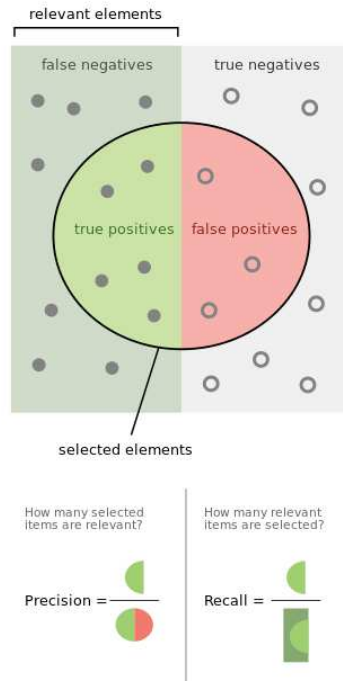


Figura 2.5: Schematizzazione visiva di TP, FP, TN, FN e delle metriche di precisione e richiamo [17].

Un metodo molto utile per visualizzare queste metriche è la matrice di confusione. Si tratta di uno strumento che mostra la frequenza di errori di classificazione e la corretta classificazione dei dati. Nelle celle di una matrice di confusione sono riportati valori numerici per indicare i vari casi (TP, FP, TN e FN). Pertanto la precisione, il richiamo e, di conseguenza, l' *F-score*, possono essere calcolati direttamente dai valori delle celle di una matrice di confusione.

In figura 2.6 è riportato uno schema d'esempio di matrice di confusione per un problema di classificazione con due classi.

		Predizioni	
		Classe P	Classe N
Realtà	Classe P	TP	FN
	Classe N	FP	TN

Figura 2.6: Struttura della matrice di confusione nel caso di un problema di classificazione con due classi: P e N .

Capitolo 3

Alberi decisionali per l'identificazione di IMBHs

Dopo aver introdotto gli strumenti teorici necessari, in questo Capitolo viene descritta nello specifico l'applicazione di tali concetti al problema astrofisico in questione: l'identificazione di IMBHs al centro di GCs. In particolare, in questo lavoro di tesi, si vuole predire la presenza di IMBHs al centro dei GCs utilizzando un metodo di ML interpretabile sulla base delle caratteristiche dinamiche delle MSPs presenti nelle regioni centrali degli ammassi. I modelli di ML utilizzati sono gli alberi decisionali costruiti con l'algoritmo CART (Cap. 2). L'interpretabilità degli alberi deriva dal fatto che a ciascun nodo la decisione presa dall'albero risulta esprimibile in termini del confronto tra una sola caratteristica (ad es. l'accelerazione della quarta pulsar nel campione) e un valore soglia: pertanto si possono ricondurre le decisioni prese, almeno nei primi nodi, ad una possibile interpretazione fisica.

In primo luogo verrà brevemente descritto il codice HiGPUUs utilizzato

per le simulazioni a N-corpi da cui provengono i *dataset* forniti all'albero decisionale. Successivamente verranno, quindi, presentati i *dataset* e poi si passerà ad una descrizione generale del codice sviluppato per la predizione.

3.1 Motivazioni per l'uso degli alberi decisionali

Nel Capitolo 1 si è parlato dell'importanza di identificare IMBHs al centro di ammassi e delle problematiche che si riscontrano utilizzando i metodi diretti. Abbiamo visto che esse sono legate principalmente alle difficoltà osservative, in quanto misurare le dispersioni di velocità all'interno della sfera di influenza del GC richiede elevatissime risoluzioni spaziali. Inoltre, l'identificazione di una cuspidi nei profili di densità e di dispersione di velocità non sempre è indice della presenza di un IMBH. Per lo sviluppo della tesi, invece, ci si è basati su un metodo di identificazione indiretto basato sulle MSPs [3]. Dallo studio degli effetti dinamici di un IMBH sulle osservabili delle MSPs all'interno dei GCs, infatti, è possibile in condizioni opportune identificare indirettamente la presenza di IMBHs.

In questo lavoro di tesi, però, si vuole applicare un nuovo metodo in questo contesto, ovvero servirsi di un algoritmo di ML interpretabile (gli alberi decisionali) per prevedere la presenza di IMBHs al centro di GCs. La sfida, infatti, sta proprio nel voler mettere a punto un metodo identificativo nuovo per cercare di avanzare nella risoluzione di un problema astrofisico ancora così aperto. Le informazioni delle MSPs che si vogliono fornire all'algoritmo per la predizione di IMBHs sono l'accelerazione e le sue derivate rispetto al tempo: *jerks* e *snaps* (*o jounce*). Queste quantità sono legate alle derivate dei periodi delle MSPs nell'ammasso e, di conseguenza, sarebbe possibile stimarle a partire da tali misure. Le campagne osservative, tuttavia, richiedono tempistiche di diversi anni per ottenere

jerks e *snap*s con una precisione adeguata, oltre a richiedere una certa regolarità e puntualità nelle osservazioni. Pertanto, i dati osservativi necessari non sono ancora a disposizione e, per questo lavoro, ci si è serviti dei dati provenienti da un set di simulazioni a N-corpi. In ogni caso, però il *training* va eseguito su dati provenienti da simulazioni perché anche avendo dati di *jerks* e *snap*s in ammassi reali non sapremmo quali fra essi ospitino o meno un IMBH. Inoltre, avere dati di simulazioni in tal senso è di fondamentale importanza per un futuro confronto tra dati osservati e dati teorici.

3.2 Simulazioni con il codice N-body diretto HiGPUs

Le simulazioni a N-corpi da cui derivano i dati forniti all'algoritmo di ML per la previsione di IMBHs nei GCs sono state ottenute mediante la nuova versione del codice *HermiteIntegratorGPUs* (*HiGPUs*) [15]. Si tratta di un codice ad alta precisione, adatto a simulare l'evoluzione nel tempo di sistemi di masse puntiformi che interagiscono tramite la forza newtoniana classica.

Il codice HiGPUs (disponibile gratuitamente in [21]) implementa un metodo di integrazione di Hermite fino al 6° ordine che calcola e utilizza accelerazioni, *jerks* e *snap*s per far avanzare le posizioni e le velocità delle stelle nel tempo. I dati relativi alle derivate dell'accelerazione su cui andremo ad effettuare il training sono quindi già disponibili in una tipica simulazione, senza che sia necessario apportare alcuna modifica al codice.

Lo schema che segue è diviso in tre fasi: predizione, valutazione e correzione. Di seguito verranno indicate le derivate dell'accelerazione rispetto al tempo con la notazione con i punti e chiameremo indistintamente gli elementi *stelle* o *particelle*. Consideriamo, quindi, un sistema composto

da N stelle e assumiamo che l' i -esima particella al tempo iniziale $t_{c,0}$ sia caratterizzata da posizione $\mathbf{r}_{i,0}$, velocità $\mathbf{v}_{i,0}$, accelerazione $\mathbf{a}_{i,0}$, *jerk* $\dot{\mathbf{a}}_{i,0}$, *snap* $\ddot{\mathbf{a}}_{i,0}$, *crackle* $\dddot{\mathbf{a}}_{i,0}$ ed un *time step* $\Delta t_{i,0}$. Inoltre, chiamiamo m il numero di particelle appartenenti allo stesso blocco temporale e che devono essere evolute allo stesso tempo $t_{c,0} + \Delta t_{i,0}$. I passi da seguire sono, dunque, i seguenti:

- *Predizione*: consiste nel calcolo di posizione, velocità e accelerazione di tutte le stelle a partire dai valori iniziali:

$$\begin{aligned}\mathbf{r}_{i,pred} &= \mathbf{r}_{i,0} + \mathbf{v}_{i,0}\Delta t_{i,0} + \frac{1}{2}\mathbf{a}_{i,0}\Delta t_{i,0}^2 + \frac{1}{6}\dot{\mathbf{a}}_{i,0}\Delta t_{i,0}^3 + \\ &\quad + \frac{1}{24}\ddot{\mathbf{a}}_{i,0}\Delta t_{i,0}^4 + \frac{1}{120}\dddot{\mathbf{a}}_{i,0}\Delta t_{i,0}^5 \\ \mathbf{v}_{i,pred} &= \mathbf{v}_{i,0} + \mathbf{a}_{i,0}\Delta t_{i,0} + \frac{1}{2}\dot{\mathbf{a}}_{i,0}\Delta t_{i,0}^2 + \frac{1}{6}\ddot{\mathbf{a}}_{i,0}\Delta t_{i,0}^3 + \\ &\quad + \frac{1}{24}\ddot{\mathbf{a}}_{i,0}\Delta t_{i,0}^4 \\ \mathbf{a}_{i,pred} &= \mathbf{a}_{i,0} + \dot{\mathbf{a}}_{i,0}\Delta t_{i,0} + \frac{1}{2}\ddot{\mathbf{a}}_{i,0}\Delta t_{i,0}^2 + \frac{1}{6}\dddot{\mathbf{a}}_{i,0}\Delta t_{i,0}^3\end{aligned}\tag{3.1}$$

- *Valutazione*: calcolo dell'accelerazione e delle sue derivate prima e seconda per tutte le particelle $m < N$ sulla base delle posizioni e velocità predette. Le mutue interazioni tra l' i -esima particella e le $N - 1$ rimanenti sono descritte dalle seguenti relazioni:

$$\begin{aligned}\mathbf{a}_{i,1} &= \sum_{\substack{j=1 \\ j \neq i}}^N \mathbf{a}_{ij,1} = \sum_{\substack{j=1 \\ j \neq i}}^N m_j \frac{\mathbf{r}_{ij}}{r_{ij}^3} \\ \dot{\mathbf{a}}_{i,1} &= \sum_{\substack{j=1 \\ j \neq i}}^N \dot{\mathbf{a}}_{ij,1} = \sum_{\substack{j=1 \\ j \neq i}}^N \left(m_j \frac{\mathbf{v}_{ij}}{r_{ij}^3} - 3\alpha_{ij}\mathbf{a}_{ij,1} \right) \\ \ddot{\mathbf{a}}_{i,1} &= \sum_{\substack{j=1 \\ j \neq i}}^N \ddot{\mathbf{a}}_{ij,1} = \sum_{\substack{j=1 \\ j \neq i}}^N \left(m_j \frac{\mathbf{a}_{ij}}{r_{ij}^3} - 6\alpha\dot{\mathbf{a}}_{ij,1} - 3\beta_{ij}\mathbf{a}_{ij,1} \right)\end{aligned}\tag{3.2}$$

dove $\mathbf{r}_{ij} \equiv \mathbf{r}_{j,pred} - \mathbf{r}_{i,pred}$, $\mathbf{v}_{ij} \equiv \mathbf{v}_{j,pred} - \mathbf{v}_{i,pred}$, $\mathbf{a}_{ij} \equiv \mathbf{a}_{j,pred} - \mathbf{a}_{i,pred}$, $\alpha_{ij}r_{ij}^2 \equiv \mathbf{r}_{ij} \cdot \mathbf{v}_{ij}$, $\beta_{ij}r_{ij}^2 \equiv v_{ij}^2 + \mathbf{r}_{ij} \cdot \mathbf{a}_{ij} + \alpha_{ij}^2r_{ij}^2$

- *Correzione*: le posizioni e le velocità delle m particelle vengono corrette per mezzo delle accelerazioni e relative derivate secondo le relazioni:

$$\begin{aligned}
 \mathbf{v}_{i,\text{corr}} &= \mathbf{v}_{i,0} + \frac{\Delta t_{i,0}}{2} (\mathbf{a}_{i,1} + \mathbf{a}_{i,0}) - \frac{\Delta t_{i,0}^2}{10} (\dot{\mathbf{a}}_{i,1} - \dot{\mathbf{a}}_{i,0}) + \\
 &\quad + \frac{\Delta t_{i,0}^3}{120} (\ddot{\mathbf{a}}_{i,1} + \ddot{\mathbf{a}}_{i,0}) \\
 \mathbf{r}_{i,\text{corr}} &= \mathbf{r}_{i,0} + \frac{\Delta t_{i,0}}{2} (\mathbf{v}_{i,\text{corr}} + \mathbf{v}_{i,0}) - \frac{\Delta t_{i,0}^2}{10} (\mathbf{a}_{i,1} - \mathbf{a}_{i,0}) + \\
 &\quad + \frac{\Delta t_{i,0}^3}{120} (\dot{\mathbf{a}}_{i,1} + \dot{\mathbf{a}}_{i,0})
 \end{aligned} \tag{3.3}$$

HiGPUs è un codice ad alta precisione perché oltre ad utilizzare l'integrazione di Hermite fino al 6° ordine, calcola sempre le interazioni mutue tra tutte le stelle e definisce un *time step* variabile per ciascuna stella sulla base delle loro accelerazioni. Questo garantisce un buon compromesso tra precisione e velocità di calcolo.

Tuttavia, utilizzando intervalli temporali differenti durante l'integrazione, nel caso di incontri ravvicinati tra stelle, potrebbe accadere che l'intervallo temporale diventi così piccolo da bloccare in effetti l'integrazione. E' il caso in cui si voglia simulare un ammasso stellare per la tipica vita media di uno di questi oggetti, dell'ordine dell'età dell'universo, in cui si formano binarie dinamicamente importanti con periodi dell'ordine di un anno. Per ovviare al problema si può utilizzare il *softening* ϵ : una costante che fa in modo che le forze di interazione ed i rispettivi potenziali vengano smorzati per piccoli $r_{i,j}$. In questo modo, il potenziale di interazione tra due particelle di massa m_i ed m_j diventerà:

$$U_{ij} = \frac{Gm_i m_j}{\sqrt{r_{ij}^2 + \epsilon^2}} \tag{3.4}$$

Per le simulazioni da cui provengono i dati forniti al modello previsionale è stato utilizzato $\epsilon = 10^{-5} pc$, un valore che determina un compromesso tra risoluzione a piccole scale e velocità computazionale delle simulazioni.

3.3 Simulazioni e condizioni iniziali

Le simulazioni da cui provengono i *dataset* utilizzati, sono state effettuate sul supercomputer MARCONI del CINECA [16] mediante l'utilizzo del codice HiGPUs.

E' stato fatto evolvere un set di ammassi stellari; nello specifico sono state utilizzate 196 simulazioni. Ogni ammasso è composto da un numero di stelle variabile fino a $N = 100000$ ed è rappresentato da una funzione di distribuzione di Plummer [39]. Questo modello di ammasso ha simmetria sferica tanto nello spazio delle posizioni che in quello delle velocità ed è caratterizzato dalla distribuzione di densità:

$$\rho(r) = \frac{3M}{4\pi a^3} \left(1 + \frac{r^2}{a^2}\right)^{-\frac{5}{2}} \quad (3.5)$$

e potenziale corrispondente:

$$\Phi_P(r) = -\frac{GM}{\sqrt{r^2 + a^2}} \quad (3.6)$$

con G costante di gravitazione universale, M massa totale dell'ammasso data dal prodotto della massa media delle particelle per il numero N di particelle, r la coordinata radiale ed a lunghezza di scala che è stata scelta circa pari a $1pc$.

Le singole particelle sono state generate seguendo una funzione di massa iniziale (*Initial Mass Function*, *IMF*) di Kroupa 2001 [24] definita come:

$$N(m)dm \propto m^{-\alpha}dm \quad (3.7)$$

dove

$$\alpha = \begin{cases} 1.3 & \text{se } 0.08 < m < 0.5M_{\odot} \\ 2.3 & \text{se } m > 0.5M_{\odot} \end{cases} \quad (3.8)$$

Le masse delle singole stelle variano tra un minimo fissato a $0.1M_{\odot}$ ed un massimo variabile nell'intervallo $10 - 150M_{\odot}$. Si è scelto di far variare

il massimo per poter esplorare i diversi effetti dinamici che esso determinerebbe su un IMBH. Dalla IMF di Kroupa, quindi, dipende la massa media delle particelle. Anche la massa centrale dell'IMBH è stata fatta variare, in questo caso tra 0 e $10^4 M_{\odot}$. Inoltre, per ragioni tecniche, ogni sistema è stato simulato fino ad un'età variabile, ma ognuno almeno fino ad un tempo di rilassamento. Poi, per costruire i *dataset* utilizzati, sono stati selezionati casualmente e secondo una distribuzione uniforme, circa 10 *snapshot* per ogni ammasso.

Il tempo di rilassamento può essere definito come il tempo necessario ad un corpo per perdere la memoria della sua velocità iniziale mediante interazioni a due corpi, ovvero il tempo entro il quale $\left| \frac{\Delta v}{v} \right| \approx 1$. Su questo tempo scala gli effetti dinamici legati alle interazioni collisionali a due corpi divengono evidenti, modificando sensibilmente la distribuzione iniziale. Infine, le simulazioni non includono nè binarie primordiali nè evoluzione stellare.

Di seguito vengono mostrati i risultati per uno degli ammassi simulati con queste condizioni iniziali.

In figura 3.1 vi è la proiezione sul piano del cielo dell'ammasso. In figura

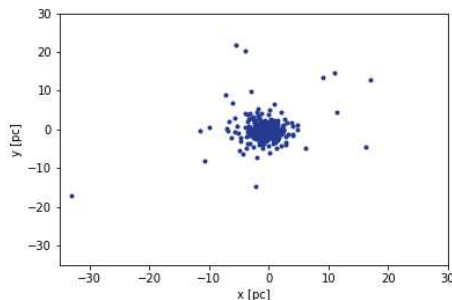
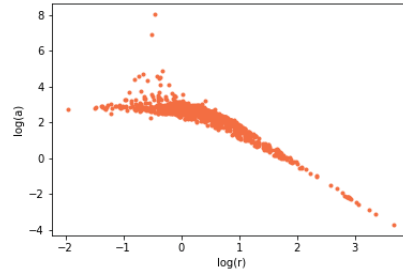


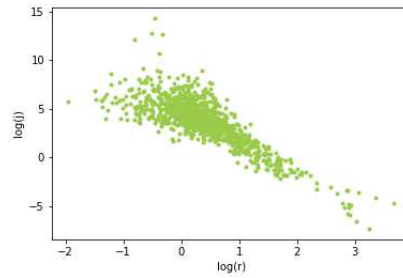
Figura 3.1: Proiezione di un ammasso campione sul piano del cielo.

3.2 sono riportati gli andamenti dell'accelerazione, del *jerk* e dello *snap* in funzione del raggio dell'ammasso. Per questi ultimi, si nota nelle regioni

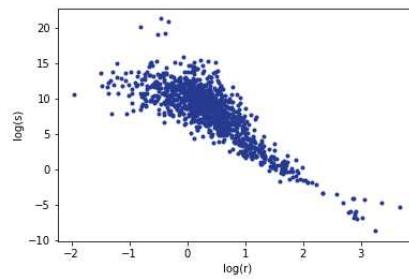
centrali una dispersione: probabilmente si tratta di particelle interagenti, in quanto presentano valori di accelerazione, *jerk* e *snap* maggiori rispetto agli andamenti generali.



(a) accelerazione VS raggio



(b) *jerk* VS raggio



(c) *snap* VS raggio

Figura 3.2: Andamenti di accelerazione, *jerk* e *snap* in funzione del raggio delle particelle nell'ammasso campione in scala logaritmica.

In figura 3.3, invece, vi sono i *jerk*s e gli *snap*s in funzione delle accelerazioni. Infine in figura 3.4 è riportato il profilo radiale di densità di massa dell'ammasso. Tutti i grafici sono in scala logaritmica.

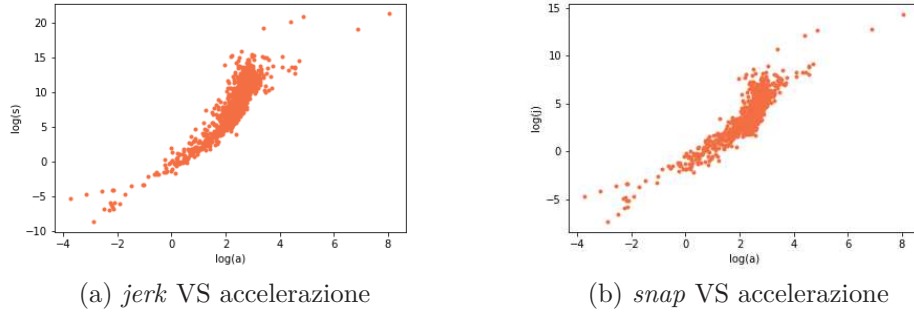


Figura 3.3: *Jerk* e *snap* in funzione dell'accelerazione delle stelle dell'ammasso in scala logaritmica.

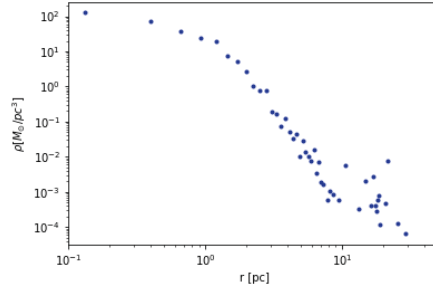


Figura 3.4: Profilo di densità dell'ammasso in scala logaritmica.

3.4 Dataset

Tra le particelle sono state etichettate come pulsar solo quelle delle zone più interne degli ammassi, in quanto, normalmente le MSPs si trovano

proprio in tali regioni a causa della massa più elevata rispetto alla stella media, che provoca segregazione.

Si è deciso di definire le zone centrali degli ammassi per mezzo del raggio che contiene (in 3D) il 20% delle stelle totali. Successivamente, in queste regioni, sono state eseguite 6 estrazioni casuali per ogni ammasso di un numero di MSPs variabile tra 1 e 40. Trattandosi di simulazioni che non tengono conto dell'evoluzione stellare, le particelle risultano tutte uguali (sono tutte dei meri tracers del campo gravitazionale), per questo motivo la scelta di estrarle casualmente è ragionevole. Tuttavia, si è scelto di selezionarle nelle regioni centrali poiché è lì che normalmente si trovano.

I *dataset* forniti all'algoritmo di ML sono in totale sei ed essi differiscono tra loro in base al diverso numero di MSPs estratte per ciascun ammasso, ovvero 1, 5, 10, 20, 30 o 40. Pertanto, ogni *record* dei diversi *dataset* corrisponde ad un ammasso e, come già detto, gli ammassi simulati sono 196. Ciascun ammasso può contenere o meno un IMBH ed è caratterizzato da un certo numero di particelle da 1 a 40 che sono state estratte casualmente come MSPs.

Ogni MSP è caratterizzata da posizione nell'ammasso, velocità, accelerazione, *jerk* e *snap*. Ma le caratteristiche delle MSPs che si è deciso di fornire come *features* all'albero decisionale sono le seguenti quantità prese in modulo:

- il raggio proiettato sul piano del cielo, cioè la distanza di ogni particella dal centro di massa dell'ammasso: $r_p = \sqrt{x^2 + y^2}$;
- la componente lungo la linea di vista (z) della velocità, v_z ;
- la componente z dell'accelerazione, a_z ;
- la componente z del *jerk*, che ora indicheremo con j_z ;
- la componente z dello *snap*, che indicheremo con s_z .

Sulla base di tali *features* per ciascun *dataset*, l'algoritmo di ML sarà in grado di predire la presenza di un IMBH negli ammassi con una certa accuratezza ed *F-score*.

3.5 Codice

Il codice sviluppato per la predizione degli IMBHs al centro di GCs mediante l'utilizzo di alberi decisionali è stato scritto in Python 3.7.0 utilizzando i moduli **NumPy** [36] e **scikit-learn** [44]. Il codice è consultabile al seguente link:

<https://gitlab.com/erica.greco/hunting-for-imbhs-with-pulsars>.

Il codice si sviluppa principalmente in quattro parti:

- Fase di preparazione del *dataset*;
- Fase di addestramento sul *training set*;
- Fase di validazione;
- Valutazione delle prestazioni sul *test set*.

Durante la prima fase, il *dataset* viene diviso nei tre sottogruppi: *training set*, *validation set* e *test set* (Sez. 2.10). Per questo lavoro si è scelto di considerare nel *test set* il 20% dei dati e il restante 80% viene suddiviso nel seguente modo: il 30% di questi dati costituisce il *validation set* ed il rimanente 70% fa parte del *training set*. Pertanto, il *dataset* è diviso in: 56% in *training set*, 24% in *validation set* e 20% in *test set*.

Nella fase successiva, l'albero decisionale di classificazione viene addestrato sul *training set*. Trattandosi di un algoritmo supervisionato, quindi, gli viene fornito in input anche il vettore delle *labels* in cui le due classi considerate vengono indicate con:

- **classe 0 o "no"**: "non c'è un IMBH nell'ammasso";

- **classe 1 o "yes"**: "c'è un IMBH nell'ammasso".

Per definire le classi si è fissato discrezionalmente un valore di soglia per la massa dell'IMBH. In particolare, si è deciso di inserire nella **classe "yes"** tutti gli ammassi per cui

$$\frac{M_{BH}}{M_{clus}} > 0.14$$

dove con M_{BH} si indica la massa del buco nero centrale e con M_{clus} la massa dell'intero ammasso. Quando, invece, la massa del buco nero centrale risulta inferiore o uguale al 14% della massa dell'intero ammasso, consideriamo che non ci sia un IMBH al centro dell'ammasso in questione. La scelta di questo valore soglia (invece di altri più comuni in letteratura come per es. il 10%) ci ha consentito di ottenere un dataset bilanciato con circa la metà degli ammassi simulati appartenenti alla **classe "yes"**.

Durante la fase di validazione, poi, avviene il processo di *pruning* dell'albero (Sez. 2.9). In questa fase si vuole scegliere il parametro di complessità α relativo al modello allenato per cui non si abbia *overfitting*, ma che permetta di migliorare l'accuratezza delle previsioni. Come già visto in maggiore dettaglio nelle Sezioni 2.9 e 2.10, per ogni valore del parametro α che l'algoritmo fa variare in maniera crescente tra 0 e 1, esiste un'unica sottostruttura dell'albero T che minimizzi la 2.2. Ad ognuno degli α selezionati e salvati in memoria dall'algoritmo corrisponde un modello allenato. Si vuole ora scegliere il migliore tra essi in modo da ottenere un classificatore in grado di predire in maniera accurata la presenza degli IMBHs al centro dei GCs. Per fare ciò si utilizza il *validation set*. In particolare, in seguito alla fase di addestramento, il codice identifica il modello classificatore corrispondente al relativo α in grado di fare una predizione sul *validation set* con la miglior accuratezza possibile.

In figura 3.5 si riporta l'andamento delle accuratezze dei modelli allenati per ciascun valore del parametro di complessità α in funzione del numero progressivo degli α stessi. Questo grafico è un esempio ottenuto a partire

dal *dataset* con l'estrazione di 20 MSPs per ogni ammasso, ma si è svolto lo stesso procedimento per tutti i *dataset*.

Il modello scelto, quindi, è quello per cui si ha il massimo valore di accuratezza nelle previsioni sul *validation set*.

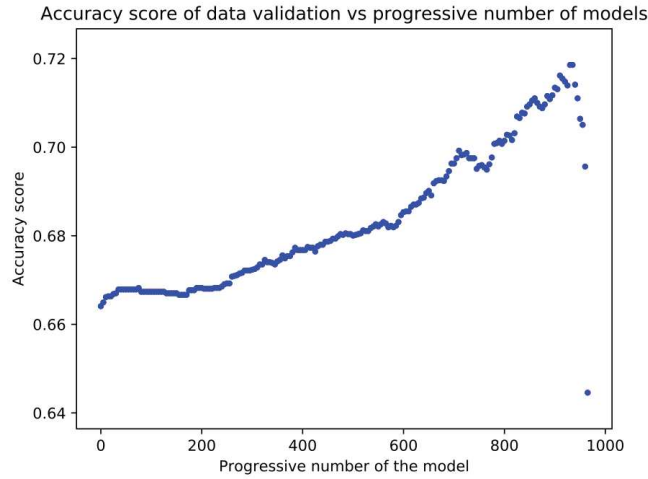


Figura 3.5: Il plot è stato ottenuto a partire dal *dataset* contenente 20 MSPs per ogni ammasso. Si riporta l'andamento dell'accuratezza ottenuta con ciascun modello allenato in funzione del numero progressivo di α : cresce fino a raggiungere il suo valore massimo con il modello relativo al 930° valore di α considerato, dopodiché diminuisce rapidamente. Il modello scelto è quello che restituisce il massimo valore di accuratezza sul *validation set*.

Una volta scelto il modello, le sue prestazioni vengono valutate su un set di dati totalmente nuovi per lui: il *test set*. Dopo il *pruning*, infatti, il modello è in grado di classificare nuovi ammassi con una certa accuratezza senza produrre *overfitting*.

Capitolo 4

Risultati

Lo scopo principale è, se vogliamo, la sfida di ogni problema di ML consiste nell'ottenere delle buone *performance* dell'algoritmo su dati nuovi e mai visti. Come descritto nel capitolo precedente, dopo aver scelto il miglior modello di albero decisionale mediante la tecnica del *pruning*, bisogna valutare le sue prestazioni sul *test set*. In questo Capitolo, dunque, vengono presentati i risultati ottenuti dal classificatore su ognuno dei *dataset* considerati. Gli strumenti utilizzati per la valutazione sono le matrici di confusione (Sez. 2.11.2) e le metriche di precisione, richiamo ed *F-score*. Infine, visualizzando l'albero è possibile analizzarne le decisioni nodo per nodo (soprattutto i primi nodi) [11] e valutare l'*importanza* delle *features*. Ciò permette di fornire un'interpretazione fisica dei risultati.

4.1 Valutazione delle performance

Sono stati utilizzati sei *dataset* diversi per la previsione della presenza di IMBHs al centro di GCs a partire dalle caratteristiche dinamiche delle MSPs che si trovano nelle regioni centrali degli ammassi. I *dataset* differiscono tra loro per il numero di MSPs selezionate casualmente in queste

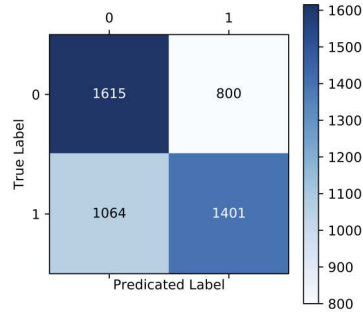
zone per ogni ammasso: 1, 5, 10, 20, 30 o 40.

Le prestazioni degli alberi decisionali vengono valutate sui relativi *test set* in base alla loro capacità nel classificare gli ammassi nella **classe 0** o nella **classe 1**.

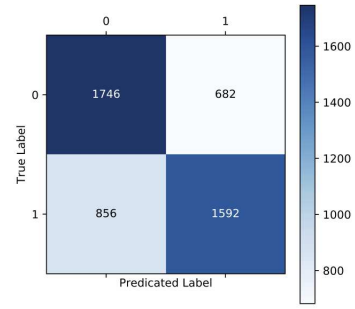
In figura 4.1 sono riportate le matrici di confusione ottenute considerando il miglior modello classificatore per ogni *dataset*. Esse sono un utile strumento per visualizzare la qualità delle classificazioni. Infatti, mettono in relazione quelli che sono i reali valori delle *labels* ("*True label*") con i valori che, invece, sono stati predetti ("*Predicted label*") per ciascuna classe. Nelle celle sulle diagonalvi è il numero di risposte corrette date dal classificatore, mentre nelle altre celle vi è il numero di risposte sbagliate. Si vede che, con l'aumentare del numero di MSPs selezionate nel centro di ciascun ammasso, aumenta anche il numero di risposte corrette da parte dell'albero decisionale. In particolare, i corrispondenti valori delle metriche di precisione, richiamo ed *F-score* sono riportati nelle tabelle 4.1. Per maggiore chiarezza, infine, in figura 4.2, vi è l'andamento della *F-score* per ciascuna classe, in funzione del numero di MSPs considerato per ogni ammasso. Si vede che, per entrambe le classi, i valori di *F-score* aumentano notevolmente (di circa il 10%) con l'aumentare del numero di MSPs, fino a 10 MSPs per ammasso. Da 20 MSPs in poi, le *F-score* si stabilizzano al 72% e 73% rispettivamente per le classi 0 e 1. Questo potrebbe essere dovuto o al fatto che l'albero non riesce ad estrarre ulteriori informazioni dal *dataset* nonostante l'incremento del numero di MSPs considerate oppure che realmente l'aumento di MSPs non fornisce ulteriori informazioni. Al fine di aumentare le prestazioni, quindi, sarebbe utile fornire all'algoritmo dei *dataset* più consistenti in termini di numero totale di ammassi, oppure cambiare algoritmo sia pure a prezzo di perdere interpretabilità, usando ad esempio una foresta casuale (la media di un insieme di alberi).

In ogni caso, considerando la natura preliminare di questo studio, si tratta di buoni risultati che potrebbero permettere, all'arrivo di dati reali, di

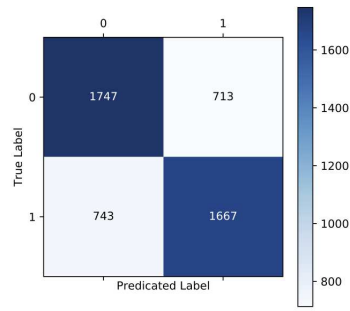
avere una buona indicazione sulla presenza o meno di IMBHs nei GCs, soprattutto in vista di studi più dettagliati e costosi in termini di tempo di telescopio, mirati quindi al singolo ammasso.



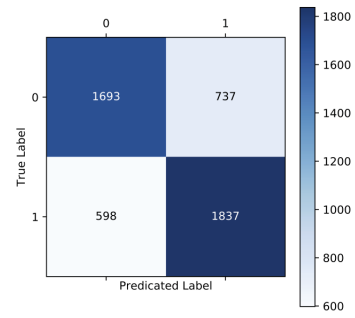
(a) 1 MSP per ammasso.



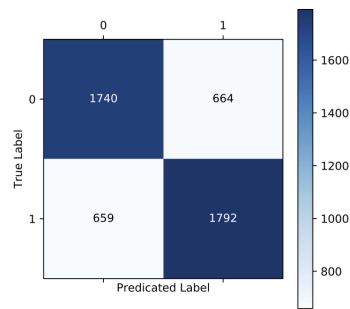
(b) 5 MSPs per ammasso.



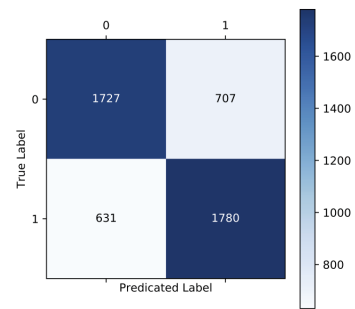
(c) 10 MSPs per ammasso.



(d) 20 MSPs per ammasso.



(e) 30 MSPs per ammasso.



(f) 40 MSPs per ammasso.

Figura 4.1: Matrici di confusione ottenute sui *test set* per i sei diversi *dataset*.

classe	precisione	richiamo	F-score
0	0.60	0.67	0.63
1	0.64	0.57	0.60

(a) 1 MSP per ammasso.

classe	precisione	richiamo	F-score
0	0.67	0.72	0.69
1	0.70	0.65	0.67

(b) 5 MSPs per ammasso.

classe	precisione	richiamo	F-score
0	0.70	0.71	0.71
1	0.70	0.69	0.70

(c) 10 MSPs per ammasso.

classe	precisione	richiamo	F-score
0	0.74	0.70	0.72
1	0.71	0.75	0.73

(d) 20 MSPs per ammasso.

classe	precisione	richiamo	F-score
0	0.73	0.72	0.72
1	0.73	0.73	0.73

(e) 30 MSPs per ammasso.

classe	precisione	richiamo	F-score
0	0.73	0.71	0.72
1	0.72	0.74	0.73

(f) 40 MSPs per ammasso.

Tabella 4.1: Risultati di precisione, richiamo ed F-score per ciascuna classe, ottenuti sui *test set* per i sei diversi *dataset*.

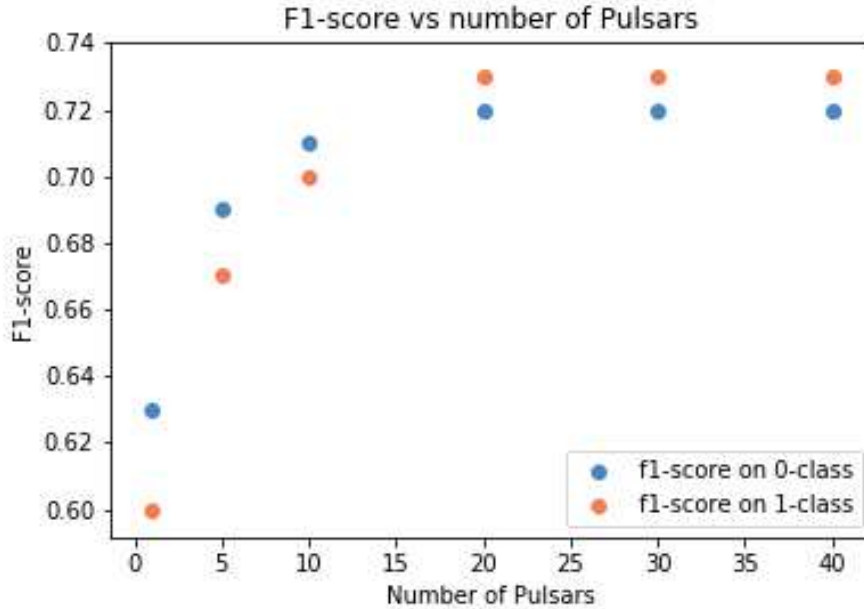


Figura 4.2: Andamento della F -score per le due distinte classi in funzione del numero di MSPs considerate per ogni ammasso.

4.2 Interpretazione dei risultati

Gli alberi decisionali, soprattutto se non sono molto profondi, sono semplici da interpretare, in quanto essi possono essere visualizzati.

In figura 4.3 è riportato un esempio di albero decisionale ottenuto con il *dataset* relativo ad una MSP estratta per ogni ammasso.

L'interpretazione è semplice: partendo dal nodo radice si percorre l'albero passando per i nodi successivi per mezzo dei rami. Questi ultimi suggeriscono quale sottoinsieme andremo a guardare. Abbiamo già visto che gli *split*, per ogni sottoinsieme, vengono decisi in base ad una misura di impurità fatta sulla valutazione dell'indice di Gini.

L'algoritmo sceglie come nodo radice quello con impurità più bassa, ossia

quello con il valore di indice di Gini minore.

Una volta deciso il nodo radice, l'albero viene ingrandito ad una profondità pari ad uno. Lo stesso processo viene ripetuto per gli altri nodi nell'albero fino a raggiungere la massima profondità. In questo caso, stiamo interpretando un albero con profondità pari a 7, ottenuta in seguito alla fase di potatura.

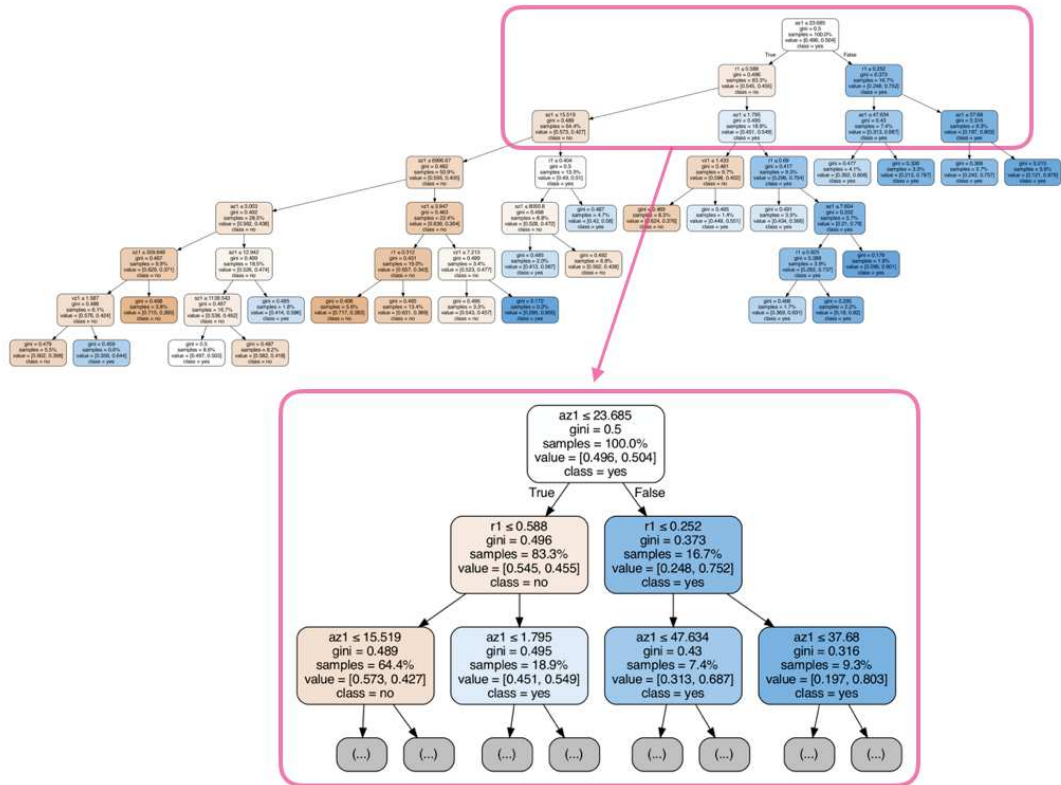


Figura 4.3: Un esempio di albero decisionale ottenuto a partire dal *dataset* di 1 MSP selezionata nella regione centrale di ogni ammasso considerato. Per chiarezza è stato eseguito uno zoom nei primi *split* dell'albero.

Per maggior chiarezza, in figura 4.3 è stato fatto uno zoom sui primi *split*, in particolare per il primo nodo abbiamo che:

- **$a_{z1} \leq 23.685$** : il primo *split* viene fatto sulla *feature* a_{z1} , ovvero ci si chiede se la componente z dell’accelerazione è minore o uguale di 23.685 e, in base al risultato, si segue il percorso *True* (freccia di sinistra) o *False* (freccia di destra). In particolare, seguendo il ramo *True* si arriva in un sottoinsieme di ammassi per il quale si prevede con maggior probabilità la **classe "no"**, ovvero quella per cui non è presente un IMBH. Ciò può essere interpretato con il fatto che valori minori di accelerazione delle MSPs siano riconducibili ad una minore attrazione gravitazionale verso le regioni centrali e, quindi, ad una massa centrale minore;
- **gini = 0.5**: è l’indice di Gini per la *feature* a_{z1} . Un valore dell’indice di Gini pari a 0 significa che il nodo è puro, ossia che esiste solo una classe (quindi che gli ammassi appartengono tutti alla **classe "yes"** o alla **classe "no"**). Al contrario, un valore maggiore di 0, come nel nodo radice, ci indica che i campioni all’interno del nodo appartengono ad entrambe le classi;
- **samples = 100%**: indica la percentuale di ammassi da classificare appartenenti al *training set*. Nel nodo radice è corretto aspettarsi che ancora tutti gli ammassi devono essere classificati;
- **value = [0.496, 0.504]**: indica la frazione di ammassi di quel nodo che ricadono in ciascuna classe. Ciò significa che nel nodo radice il 49.6% degli ammassi appartiene alla **classe "no"** (che abbiamo anche chiamato **classe 0**) ed il restante 50.4% fa parte della **classe "yes"** (chiamata anche **classe 1**);
- **class = yes**: rappresenta la previsione su quel determinato nodo e si ricava da **value**. Infatti, dal punto sopra, capiamo che nel nodo c’è una percentuale maggiore di campioni appartenenti alla **classe "yes"**.

Queste denominazioni vengono ripetute per i vari nodi successivi dell'albero. Nello specifico, si può notare che la somma delle *samples* dei nodi figli restituisce il valore *samples* dei nodi padri.

Inoltre, la **classe "yes"** è identificata con l'azzurro, mentre la **classe "no"** con l'arancione. L'intensità della colorazione di ogni nodo varia in base all'accuratezza della previsione: vediamo che nei nodi di azzurro più intenso una percentuale maggiore di ammassi viene classificata in **"yes"** e in quelli di arancione più intenso, invece, una percentuale maggiore di ammassi viene classificata in **"no"**.

Infine, nei nodi foglia non appaiono più le *features*, in quanto si è raggiunta la massima profondità e tutti gli ammassi in quel nodo sono stati classificati.

In ultimo, è possibile valutare l'importanza delle diverse *features* per un albero decisionale. L'importanza complessiva di una singola *feature* può essere calcolata nel modo seguente: bisogna esaminare tutti gli *split* per i quali la *feature* è stata utilizzata e misurare quanto essa ha ridotto l'indice di Gini rispetto al nodo padre. La somma di tutte le importanze viene scalata a 100, pertanto, ogni importanza può essere interpretata come una quota dell'importanza complessiva per il modello [35]. Nelle figure 4.4 e 4.5 vengono riportati gli istogrammi relativi all'importanza delle diverse *features* per ciascuno dei *dataset* utilizzati. Le MSPs sono ordinate per raggi r_p proiettati sul piano del cielo crescenti. E, le altre *features* sono i rispettivi valori del modulo di velocità v_z , accelerazione a_z , *jerk* j_z e *snap* s_z delle MSPs lungo la linea di vista. Si può notare che, in tutti i casi, le *features* con importanza maggiore sono: i raggi, in particolare quello dell'ultima MSP, le accelerazioni e gli *snap*.

Che le accelerazioni fossero caratterizzate da un'alta importanza era un risultato atteso. Una possibile spiegazione, invece, per l'importanza del raggio dell'ultima MSP potrebbe essere che il modello lo interpreti come una misura della dimensione dell'ammasso, o meglio, come una misura dello spazio in cui si trovano le MSPs considerate.

Per quanto riguarda, invece, l'interpretazione fisica della ridotta importanza delle *features* associate a derivate dispari della posizione (velocità e *jerk*) rispetto alle derivate pari (accelerazione e *snap*) non è al momento completamente chiara e meriterebbe ulteriore investigazione. In ogni caso, va notato come le pulsar selezionate, per costruzione particelle che si trovano nella zona più interna dell'ammasso, evolvano in un potenziale medio che, trascurando per ora il contributo dell'IMBH, è sostanzialmente armonico poiché il *core* di una tipica simulazione ha densità pressoché costante, pertanto:

$$V = \frac{1}{2}kr^2$$

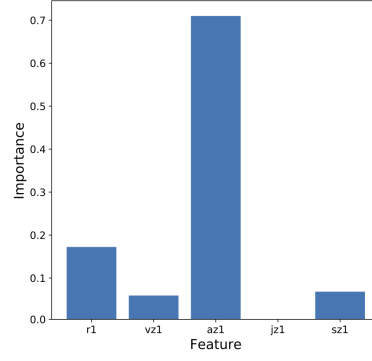
dove $k = \frac{4}{3}\pi G\rho$, con ρ densità costante del *core* ed r distanza dal centro del *core*.

Il moto armonico che ne consegue impartisce, dal momento che si disaccoppiano, a ciascuna coordinata una dipendenza nella forma:

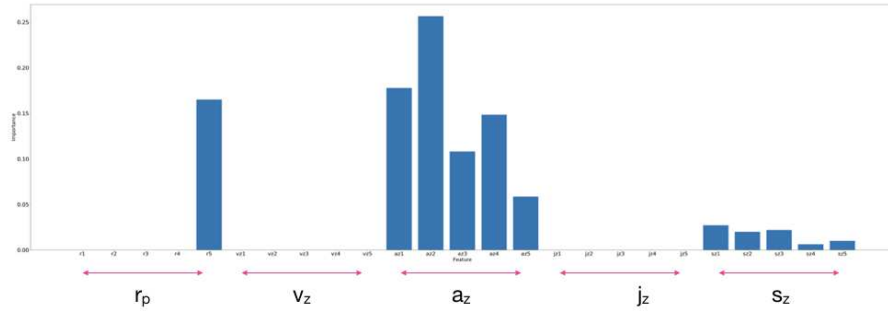
$$\begin{cases} x = x_0 \sin(\omega t + \phi_x) \\ y = y_0 \sin(\omega t + \phi_y) \\ z = z_0 \sin(\omega t + \phi_z) \end{cases}$$

le cui derivate prime sono proporzionali al coseno, le derivate seconde nuovamente al seno, le derivate terze ancora al coseno, ecc...

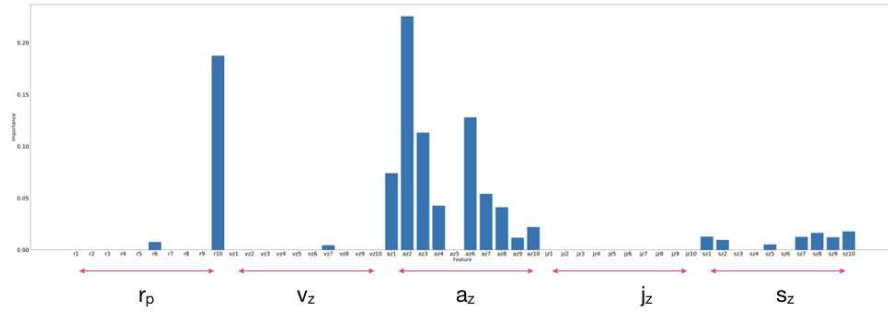
Pertanto, non è sorprendente che vi siano forti correlazioni tra le derivate dispari (velocità e *jerk*, proporzionali al coseno) e tra le derivate pari (accelerazione e *snap*, proporzionali al seno).



(a) 1 MSP per ammasso.

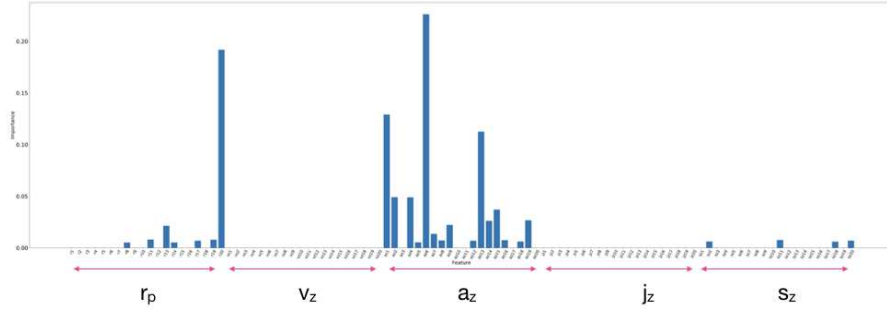


(b) 5 MSPs per ammasso.

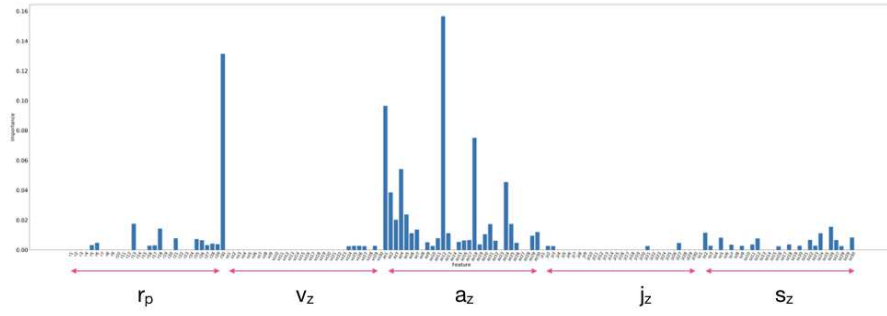


(c) 10 MSPs per ammasso.

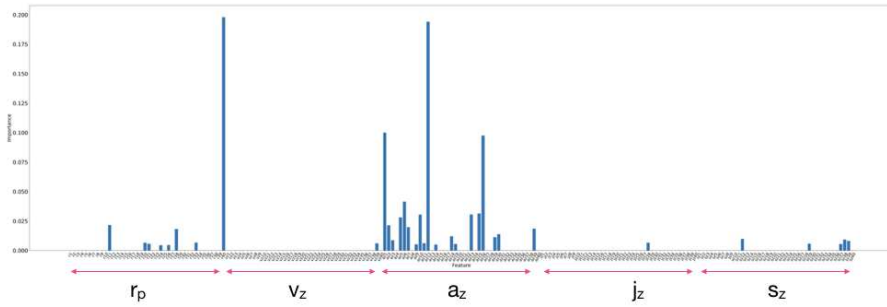
Figura 4.4: Grafici delle importanze delle *features* per il modello calcolate per i *dataset* relativi a 1, 5 e 10 MSPs per ammasso (dall'alto verso il basso).



(a) 20 MSPs per ammasso.



(b) 30 MSPs per ammasso.



(c) 40 MSPs per ammasso.

Figura 4.5: Grafici delle importanze delle *features* per il modello calcolate per i *dataset* relativi a 20, 30 e 40 MSPs per ammasso (dall'alto verso il basso).

Capitolo 5

Conclusioni e sviluppi futuri

Gli IMBHs, la loro esistenza ed il fatto di poterli trovare al centro dei GCs sono stati per lungo tempo oggetto di dibattito presso la comunità astronomica. I modelli teorici comunemente accettati (merger ripetuto di buchi neri di massa stellare oppure fusione di stelle a produrre una very massive star nell'ammasso) indicano che l'ambiente denso degli ammassi globulari sarebbe un luogo privilegiato per la formazione di IMBH. Tuttavia, i metodi osservativi usati finora non ne hanno ancora verificato la presenza in modo ritenuto definitivo, anche a causa del fatto che osservazioni dirette di emissione elettromagnetica da accrescimento sono precluse dalla quasi completa assenza di gas negli ammassi globulari. Tra i metodi indiretti, che si basano sull'osservazione degli effetti dinamici di un IMBH sulle stelle visibili dell'ammasso, un interessante metodo proposto per identificarli è quello che si basa sullo studio dinamico delle MSPs. Da qui nasce l'idea di questo lavoro di tesi: identificare gli IMBHs in ammassi globulari attraverso una tecnica di ML interpretabile utilizzando le osservabili cinematiche delle MSPs. Si tratta, infatti, di applicare

un metodo nuovo ad un problema ben noto, ma ancora sostanzialmente aperto.

A partire dai risultati di un set di simulazioni a N-corpi di ammassi globulari sono stati creati i *dataset* forniti all'algoritmo di ML. Lo scopo è quello di studiare la presenza di IMBHs in questi ammassi sulla base della posizione nel piano del cielo delle MSPs e delle loro componenti lungo la linea di vista di velocità, accelerazione, *jerk* e *snap*. Per farlo ci si è serviti degli alberi decisionali.

Sono stati considerati circa 200 ammassi, per ognuno dei quali sono state effettuate 6 estrazioni casuali di un numero variabile tra 1 e 40 di MSPs selezionate casualmente dalle regioni centrali degli ammassi stessi.

Nonostante gli alberi decisionali non siano tra i modelli di ML più performanti (sono stati infatti scelti qui per la loro semplicità e potenziale interpretabilità piuttosto che per altre caratteristiche), si ottengono dei risultati promettenti. In particolare, si riescono a classificare gli ammassi in base alla presenza di un IMBH con valori di accuratezza pari a circa il 70%.

Inoltre, la caratteristica fondamentale degli alberi decisionali è che essi possono essere facilmente interpretati, specie quando si tratta di alberi non molto profondi. Dal momento che nella scienza la capacità di spiegare in modo comprensibile le previsioni effettuate da un modello è tanto importante quanto l'accuratezza, a differenza che in altre applicazioni, l'uso di modelli di ML relativamente semplici sta riscuotendo notevole interesse in astronomia.

Dall'interpretazione delle scelte che il modello compie in ogni nodo dell'albero è possibile capire quali sono state le caratteristiche dinamiche delle MSPs che hanno avuto più *importanza* nella fase di allenamento del modello. Per tutti i casi analizzati risulta che, come atteso, l'importanza maggiore (dal linguaggio tecnico del ML "*feature importance*") viene attribuita alle accelerazioni delle MSPs negli ammassi. Molto importante

risulta anche il raggio proiettato sul piano del cielo della MSP più distante. Il resto dell'importanza è distribuito tra gli *snap*s ed i raggi proiettati delle altre MSPs.

Nella pratica, dopo aver monitorato un certo numero di MSPs in un ammasso misurandone alcune osservabili cinematiche, quello che si vuole capire è se, in base ai loro valori, l'ammasso possa ospitare un IMBH o meno. Se in questa fase avessimo già allenato un modello di ML a partire da dati di simulazioni, allora si potrebbe prevedere nella realtà la presenza di tali oggetti tanto ricercati. Pertanto, sarebbe sicuramente utile ampliare il numero di simulazioni utilizzate per creare i *dataset* forniti all'algoritmo.

L'importanza di questo lavoro sta proprio nel fatto che sarebbe possibile ottenere un riscontro futuro su dati veri. Per tale ragione, sarebbe fondamentale affinare questo metodo e svilupparne altri a partire da questo, per ottenere previsioni più accurate. Ad esempio, si potrebbe utilizzare un algoritmo più complesso, come quello delle foreste.

Un primo passo verso il miglioramento delle *performance* dell'algoritmo è quello di rendere le simulazioni da cui provengono i *dataset* il più verosimili possibile, ad esempio considerando anche l'evoluzione stellare. In secondo luogo, si potrebbero migliorare alcuni punti importanti del codice di previsione. Infatti, il criterio considerato per determinare la presenza di un IMBH in un ammasso si basa su un valore di soglia di massa relativa scelto a priori. Questo passaggio è cruciale per la classificazione degli ammassi nelle due distinte classi, per questo sarebbe importante svolgere, in futuro, uno studio più approfondito sulla scelta di tale criterio.

Bibliografia

- [1] *47 Tucanae*. URL: https://en.wikipedia.org/wiki/47_Tucanae.
- [2] Abbate F., Possenti A., Ridolfi A., Freire P. C. C., Camilo F., Manchester R. N., D'Amico N. «Internal gas models and central black hole in 47 Tucanae using millisecond pulsars». In: *MNRAS* 481.1 (nov. 2018), pp. 627–638. DOI: [10.1093/mnras/sty2298](https://doi.org/10.1093/mnras/sty2298). arXiv: [1808.06621](https://arxiv.org/abs/1808.06621) [[astro-ph.HE](#)].
- [3] Abbate F., Spera M., Colpi M. «Intermediate mass black holes in globular clusters: effects on jerks and jounces of millisecond pulsars». In: *MNRAS* 487.1 (lug. 2019), pp. 769–781. DOI: [10.1093/mnras/stz1330](https://doi.org/10.1093/mnras/stz1330). arXiv: [1905.05209](https://arxiv.org/abs/1905.05209) [[astro-ph.HE](#)].
- [4] Abbott R., Abbott T. D. et al. 2020. «Properties and Astrophysical Implications of the 150 $M_{\odot}$ Binary Black Hole Merger GW190521». In: *The Astrophysical Journal* 900.1 (2020), p. L13. DOI: [10.3847/2041-8213/aba493](https://doi.org/10.3847/2041-8213/aba493). URL: <https://doi.org/10.3847/2041-8213/aba493>.
- [5] Abel T., Bryan G. L. and Normann M. L. «The Formation and Fragmentation of Primordial Molecular Clouds». In: *ApJ* 540.1 (set. 2000), pp. 39–44. DOI: [10.1086/309295](https://doi.org/10.1086/309295). arXiv: [astro-ph/0002135](https://arxiv.org/abs/astro-ph/0002135) [[astro-ph](#)].

- [6] Abbott B. P. et al. «Observation of Gravitational Waves from a Binary Black Hole Merger». In: *Phys. Rev. Lett.* 116 (6 2016), p. 061102. DOI: [10.1103/PhysRevLett.116.061102](https://doi.org/10.1103/PhysRevLett.116.061102). URL: <https://link.aps.org/doi/10.1103/PhysRevLett.116.061102>.
- [7] Amaro-Seoane P., Gair J. R., Pound A., Hughes S. A., Sopuerta C. F. «Research Update on Extreme-Mass-Ratio Inspirals». In: *Journal of Physics Conference Series* 610, 012002 (mag. 2015), p. 012002. DOI: [10.1088/1742-6596/610/1/012002](https://doi.org/10.1088/1742-6596/610/1/012002). arXiv: [1410.0958](https://arxiv.org/abs/1410.0958) [[astro-ph.CO](#)].
- [8] *Applied Data Mining and Statistical Learning*. URL: <https://online.stat.psu.edu/stat508/>.
- [9] *Apprendimento automatico*. URL: https://it.wikipedia.org/wiki/Apprendimento_automatico.
- [10] M. Arca-Sedda e R. Capuzzo-Dolcetta. «The globular cluster migratory origin of nuclear star clusters». In: *MNRAS* 444.4 (nov. 2014), pp. 3738–3755. DOI: [10.1093/mnras/stu1683](https://doi.org/10.1093/mnras/stu1683). arXiv: [1405.7593](https://arxiv.org/abs/1405.7593) [[astro-ph.GA](#)].
- [11] Ammar Askar et al. «Finding black holes with black boxes - using machine learning to identify globular clusters with black hole sub-systems». In: *MNRAS* 485.4 (giu. 2019), pp. 5345–5362. DOI: [10.1093/mnras/stz628](https://doi.org/10.1093/mnras/stz628). arXiv: [1811.06473](https://arxiv.org/abs/1811.06473) [[astro-ph.GA](#)].
- [12] Bahcall J. N., Wolf R. A. «Star distribution around a massive black hole in a globular cluster.» In: *ApJ* 209 (ott. 1976), pp. 214–232. DOI: [10.1086/154711](https://doi.org/10.1086/154711).
- [13] Baumgardt H., Makino J., Ebisuzaki T. «Massive Black Holes in Star Clusters. II. Realistic Cluster Models». In: *ApJ* 613.2 (ott. 2004), pp. 1143–1156. DOI: [10.1086/423299](https://doi.org/10.1086/423299). arXiv: [astro-ph/0406231](https://arxiv.org/abs/astro-ph/0406231) [[astro-ph](#)].

- [14] Breiman L., Friedman J.H., Olshen R. A. and Stone C. J. *Classification and Regression Trees*. A cura di Chapman e NY Hall New York. 1984.
- [15] R. Capuzzo-Dolcetta, M. Spera e D. Punzo. «A fully parallel, high precision, N-body code running on hybrid computing platforms». In: *Journal of Computational Physics* 236 (mar. 2013), pp. 580–593. DOI: [10.1016/j.jcp.2012.11.013](https://doi.org/10.1016/j.jcp.2012.11.013). arXiv: [1207.2367](https://arxiv.org/abs/1207.2367) [[astro-ph.IM](#)].
- [16] CINECA. URL: <https://www.cineca.it>.
- [17] *F1 score*. URL: https://en.wikipedia.org/wiki/F1_score.
- [18] P. C. C. Freire et al. «Long-term observations of the pulsars in 47 Tucanae - II. Proper motions, accelerations and jerks». In: *MNRAS* 471.1 (ott. 2017), pp. 857–876. DOI: [10.1093/mnras/stx1533](https://doi.org/10.1093/mnras/stx1533). arXiv: [1706.04908](https://arxiv.org/abs/1706.04908) [[astro-ph.HE](#)].
- [19] Gebhardt K., Rich R. M., Ho L. C. «An Intermediate-Mass Black Hole in the Globular Cluster G1: Improved Significance from New Keck and Hubble Space Telescope Observations». In: *ApJ* 634.2 (dic. 2005), pp. 1093–1102. DOI: [10.1086/497023](https://doi.org/10.1086/497023). arXiv: [astro-ph/0508251](https://arxiv.org/abs/astro-ph/0508251) [[astro-ph](#)].
- [20] Martin G. Haehnelt e Martin J. Rees. «The formation of nuclei in newly formed galaxies and the evolution of the quasar population». In: *MNRAS* 263.1 (lug. 1993), pp. 168–178. DOI: [10.1093/mnras/263.1.168](https://doi.org/10.1093/mnras/263.1.168).
- [21] *HiGPUs*. URL: astrowww.phys.uniroma1.it/dolcetta/HPCcodes/HiGPUs.html.
- [22] Jenny E. Greene, Jay Strader, Luis C. Ho. «Intermediate-Mass Black Holes». In: *Annual Review of Astronomy and Astrophysics* 58.1 (2020), null. DOI: [10.1146/annurev-astro-032620-021835](https://doi.org/10.1146/annurev-astro-032620-021835). eprint: [https://doi.org/10.1146/annurev-astro-032620-](https://doi.org/10.1146/annurev-astro-032620-021835)

021835. URL: <https://doi.org/10.1146/annurev-astro-032620-021835>.
- [23] King I. «The structure of star clusters. I. an empirical density law». In: *AJ* 67 (ott. 1962), p. 471. DOI: [10.1086/108756](https://doi.org/10.1086/108756).
- [24] Pavel Kroupa. «On the variation of the initial mass function». In: *MNRAS* 322.2 (apr. 2001), pp. 231–246. DOI: [10.1046/j.1365-8711.2001.04022.x](https://doi.org/10.1046/j.1365-8711.2001.04022.x). arXiv: [astro-ph/0009005](https://arxiv.org/abs/astro-ph/0009005) [astro-ph].
- [25] Lanzoni B., Mucciarelli A., Origlia L., Bellazzini M., Ferraro F.R., Valenti E., Miocchi P., Dalessandro E., Pallanca C., Massari D. «The velocity dispersion profile of NGC 6388 from resolved-star spectroscopy: no evidence of a central cusp and new constraints on the black hole mass». In: *Astrophys. J.* 769 (2013), p. 107. DOI: [10.1088/0004-637X/769/2/107](https://doi.org/10.1088/0004-637X/769/2/107). arXiv: [1304.2953](https://arxiv.org/abs/1304.2953) [astro-ph.SR].
- [26] *LIGO-Scientific Collaboartio*. URL: <https://www.ligo.org/science/overview.php>.
- [27] Lin D. et al 2020. «Multiwavelength Follow-up of the Hyperluminous Intermediate-mass Black Hole Candidate 3XMM J215022.4-055108». In: *ApJ* 892.2, L25 (apr. 2020), p. L25. DOI: [10.3847/2041-8213/ab745b](https://doi.org/10.3847/2041-8213/ab745b). arXiv: [2002.04618](https://arxiv.org/abs/2002.04618) [astro-ph.HE].
- [28] Lin D. et al.2018. In: *Nature Astronomy* (2018), pp. 2, 656.
- [29] Lützgendorf N., Kissler-Patig M., Noyola E., Jalali B., de Zeeuw P. T., Gebhardt K., Baumgardt H. «Kinematic signature of an intermediate-mass black hole in the globular cluster NGC 6388». In: *Astronomy & Astrophysics* 533, A36 (set. 2011), A36. DOI: [10.1051/0004-6361/201116618](https://doi.org/10.1051/0004-6361/201116618). arXiv: [1107.4243](https://arxiv.org/abs/1107.4243) [astro-ph.GA].
- [30] Nora Lützgendorf et al. «Re-evaluation of the central velocity-dispersion profile in NGC 6388». In: *Astronomy & Astrophysics* 581, A1 (set. 2015), A1. DOI: [10.1051/0004-6361/201425524](https://doi.org/10.1051/0004-6361/201425524). arXiv: [1507.02813](https://arxiv.org/abs/1507.02813) [astro-ph.GA].

- [31] Mitchell Tom M. *Machine Learning*. A cura di New York: McGraw-Hill. 1997.
- [32] Maccarone T. J. «Radio emission as a test of the existence of intermediate-mass black holes in globular clusters and dwarf spheroidal galaxies». In: *MNRAS* 351.3 (lug. 2004), pp. 1049–1053. DOI: [10.1111/j.1365-2966.2004.07859.x](https://doi.org/10.1111/j.1365-2966.2004.07859.x). arXiv: [astro-ph/0403530](https://arxiv.org/abs/astro-ph/0403530) [[astro-ph](#)].
- [33] Maccarone T. J., Servillat M. «Radio observations of NGC 2808 and other globular clusters: constraints on intermediate-mass black holes». In: *MNRAS* 389.1 (set. 2008), pp. 379–384. DOI: [10.1111/j.1365-2966.2008.13577.x](https://doi.org/10.1111/j.1365-2966.2008.13577.x). arXiv: [0806.2387](https://arxiv.org/abs/0806.2387) [[astro-ph](#)].
- [34] Miller M. C., Hamilton D. P. «Production of intermediate-mass black holes in globular clusters». In: *MNRAS* 330.1 (feb. 2002), pp. 232–240. DOI: [10.1046/j.1365-8711.2002.05112.x](https://doi.org/10.1046/j.1365-8711.2002.05112.x). arXiv: [astro-ph/0106188](https://arxiv.org/abs/astro-ph/0106188) [[astro-ph](#)].
- [35] Christoph Molnar. *Interpretable Machine Learning. A Guide for Making Black Box Models Explainable*. <https://christophm.github.io/interpretable-ml-book/>. 2019.
- [36] *NumPy*. URL: <https://numpy.org>.
- [37] Pasquato M., Miocchi P., Sohn Bong W. and Young-Wook L. «Globular Clusters Hosting Intermediate-Mass Black Holes: No Mass-Segregation Based Candidates». In: *ApJ* 823.2, 135 (giu. 2016), p. 135. DOI: [10.3847/0004-637X/823/2/135](https://doi.org/10.3847/0004-637X/823/2/135). arXiv: [1604.03554](https://arxiv.org/abs/1604.03554) [[astro-ph.GA](#)].
- [38] Perera B. B. P. et al. «Evidence for an intermediate-mass black hole in the globular cluster NGC 6624». In: *MNRAS* 468.2 (giu. 2017), pp. 2114–2127. DOI: [10.1093/mnras/stx501](https://doi.org/10.1093/mnras/stx501). arXiv: [1705.01612](https://arxiv.org/abs/1705.01612) [[astro-ph.HE](#)].

- [39] H. C. Plummer. «On the problem of distribution in globular star clusters». In: *MNRAS* 71 (mar. 1911), pp. 460–470. DOI: [10.1093/mnras/71.5.460](https://doi.org/10.1093/mnras/71.5.460).
- [40] Portegies Zwart S. F. et al. 2004. In: *Nature* 428 (2004), pp. 724–726.
- [41] Portegies Zwart S. F., McMillan S. L. W. «The Runaway Growth of Intermediate-Mass Black Holes in Dense Star Clusters». In: *ApJ* 576.2 (2002), pp. 899–907. DOI: [10.1086/341798](https://doi.org/10.1086/341798). URL: <https://doi.org/10.10862F341798>.
- [42] Prager B. J., Ransom S. M., Freire P. C. C. «Using Long-term Millisecond Pulsar Timing to Obtain Physical Characteristics of the Bulge Globular Cluster Terzan 5». In: *ApJ* 845.2, 148 (ago. 2017), p. 148. DOI: [10.3847/1538-4357/aa7ed7](https://doi.org/10.3847/1538-4357/aa7ed7). arXiv: [1612.04395 \[astro-ph.SR\]](https://arxiv.org/abs/1612.04395).
- [43] *Pulsars in globular clusters*. URL: <http://www.naic.edu/%5C%7Epfreire/GCpsr.html>.
- [44] *scikit-learn*. URL: <https://www.sklearn.org>.
- [45] Shklovskii I. S. In: *Soviet Ast.* (1970), pp. 13, 562.
- [46] *Stella binaria a raggi X*. URL: https://it.wikipedia.org/wiki/Stella_binaria_a_raggi_X.
- [47] Jay Strader et al. «No Evidence for Intermediate-mass Black Holes in Globular Clusters: Strong Constraints from the JVL A». In: *ApJ* 750.2, L27 (mag. 2012), p. L27. DOI: [10.1088/2041-8205/750/2/L27](https://doi.org/10.1088/2041-8205/750/2/L27). arXiv: [1203.6352 \[astro-ph.HE\]](https://arxiv.org/abs/1203.6352).
- [48] *The Virgo Collaboration*. URL: <http://public.virgo-gw.eu/virgo-in-a-nutshell/>.

BIBLIOGRAFIA

- [49] Michele Trenti, Enrico Vesperini e Mario Pasquato. «Tidal Disruption, Global Mass Function, and Structural Parameter Evolution in Star Clusters». In: *ApJ* 708.2 (gen. 2010), pp. 1598–1610. DOI: [10.1088/0004-637X/708/2/1598](https://doi.org/10.1088/0004-637X/708/2/1598). arXiv: [0911.3394](https://arxiv.org/abs/0911.3394) [[astro-ph.GA](#)].
- [50] Volonteri M. «Formation of supermassive black holes». In: *Astronomy and Astrophysics Review* 18.3 (lug. 2010), pp. 279–315. DOI: [10.1007/s00159-010-0029-x](https://doi.org/10.1007/s00159-010-0029-x). arXiv: [1003.4404](https://arxiv.org/abs/1003.4404) [[astro-ph.CO](#)].

Ringraziamenti

Vorrei ringraziare Michela per aver creduto in questo progetto, grazie per il tuo tempo e per la disponibilità che mi hai sempre dimostrato.

Grazie di cuore Mario per avermi seguita ed accompagnata in questi mesi di lavoro, nonostante tutti gli impegni e le difficoltà dell'ultimo periodo. Con il vostro entusiasmo siete riusciti ad appassionarmi e a motivarmi. Qualsiasi cosa farò nel mio futuro voglio farla con la stessa passione che voi mettete nel vostro lavoro.

Un grazie speciale va alla mia famiglia: mamma, papà e Angelica. Sembra ieri quel giorno che iniziò quest'avventura così lontana da voi, lontana dal posto che avevo sempre chiamato Casa. Grazie per avermi sostenuta ed incoraggiata in ogni momento, il vostro ineguagliabile amore mi ha permesso di arrivare fin qui con la consapevolezza di non essere mai sola. Scusa, Angè, se ogni tanto svitavo le lampadine.

E' grazie a te, nonna, se ancora oggi continuo a ripetermi "io sono forte" prima di ogni prova. Spero che il mio Angelo preferito oggi stia lì a guardarci e a festeggiare con noi. Grazie alle mie zie, fonti di preziosi consigli, ai miei zii, ad Albi e Peppe per attendermi con ansia ad ogni mio fugace ritorno.

Alle mie amiche di sempre e a Maria Frà, va un grazie per niente banale: grazie per esserci comunque, anche a chilometri di distanza, siete sempre state una conferma, nonostante i miei innumerevoli "non posso, devo studiare". Rivedervi, per me, ogni volta è una boccata d'aria fresca:

"la felicità non è felicità senza una Capra che suona il violino".

In questi anni ho avuto la fortuna immensa di incontrare persone splendide e sarò sempre incredibilmente grata per questo.

Franci, abbiamo portato insieme un po' di Molise fino a Padova, in Islanda e alle Canarie. Sei stata il mio sostegno più grande in questi anni, grazie per esserci sempre stata. Grazie a Claudia e Raffa per essere state le prime ad accogliermi e a Laura per aver risollevato il mio animo durante il mio periodo più buio a Padova. Grazie a TEMAF e a tutti gli astronofisici per essere stati i miei primi Veri compagni di avventure e sventure. Grazie Francesco per ogni "Venga, animo!" ad ogni mio passo canario: non sarei mai arrivata in cima al possente Teide senza di te.

Non posso non dire grazie alla famiglia Astronote, alle Angurie e alla Mondanità: avete decisamente dato una svolta alla mia vita a Padova riempiendo questo cammino di suoni e sfumature che non pensavo potessero esistere.

Grazie ai così perfettamente folli astronomi del Criminal Moka Book-club (essenzialmente CMB), Giuli & Piè, Ame, Lia, Robi, Andre, Tobi, Bolli, Giusto, MartiRani e MartyParty. Il nostro prezioso legame è ciò che mi ha permesso di superare un intero lock-down e questi ultimi mesi. Grazie a voi premere quel tasto "invio" fa un po' meno paura ora.

Matteo, tutti i grazie che ho a disposizione nella mia intera vita non sarebbero mai sufficienti per ringraziarti abbastanza. Sei il mio porto sicuro. Senza di te non sarei mai riuscita ad arrivare alla fine. Anche stavolta avevi ragione tu.