



UNIVERSITÀ DEL PIEMONTE ORIENTALE

DIPARTIMENTO DI SCIENZE E INNOVAZIONE TECNOLOGICA

**Corso di Laurea Magistrale in Intelligenza Artificiale e Innovazione
Digitale**

TESI DI LAUREA

**Un approccio per la corrispondenza semantica tra insiemi di
concetti ontologici**

Relatore:

Prof. Paolo Terenziani



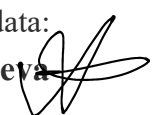
Correlatore:

Dott.ssa Erica Raina

Erica Raina

Candidata:

Anastasiia Andreeva



Matricola 20037180

Anno Accademico 2025/2026

Indice

1	Introduzione.....	4
1.1	Contesto e motivazione: CDSS e linee guida computer-interpretabili	4
1.2	Obiettivi e contributo della tesi.....	5
2	Contesto applicativo: GLARE-Edu e SNOMED CT	6
2.1	Clinical practice guidelines e pathway clinici	6
2.2	Il sistema GLARE e l'ambiente GLARE-Edu.....	6
2.3	Rappresentazione del pathway: nodi, archi e decisioni	7
2.4	SNOMED CT: concetti, identificativi e gerarchia tassonomica.....	9
2.5	Collegamento tra nodi GLARE e concetti SNOMED	10
3	Metodi per la verifica semantica	11
3.1	Similarità e distanza semantica tra concetti	11
3.2	Information Content e Lowest Common Subsumer	12
3.3	Distanza Jiang-Conrath.....	14
3.4	Adattamento dell'Information Content alla gerarchia SNOMED CT	15
3.5	Problema di assegnamento e Hungarian Algorithm	16
4	Progettazione e sviluppo del modulo	19
4.1	Dal flusso originale alla verifica semantica per segmenti	19
4.2	Estrazione delle soluzioni corrette dal segmento.....	19
4.3	Acquisizione delle risposte tramite selezione ontologica	20
4.4	Calcolo della distanza semantica	21
4.4.1	Utilizzo della transitive closure.....	21
4.4.2	Calcolo della distanza Jiang-Conrath.....	21
4.5	Matching e valutazione delle risposte.....	22
4.5.1	Valutazione dei nodi a concetto singolo	23
4.5.2	Valutazione dei nodi composti.....	23
4.6	Calcolo dello score didattico e feedback allo studente	24
4.7	Integrazione nel flusso di GLARE-Edu e ottimizzazione dell'esecuzione.....	25
5	Esempi di funzionamento e risultati	27
5.1	Impostazione e criteri di lettura degli esempi	27
5.2	Analisi del Segmento 1	28
5.2.1	Work action.....	28
5.2.2	Data enquiry	31
5.3	Analisi del Segmento 2	41

5.3.1 Work action	41
5.3.2 Data enquiry	44
5.3.3 Pharmacological prescription	48
5.4 Flusso della verifica dal punto di vista dello studente	54
5.4.1 Avvio della verifica	54
5.4.2 Svolgimento della verifica nel Segmento 1	54
5.4.3 Svolgimento della verifica nel Segmento 2	60
5.4.4 Conclusione dell'esercizio e riepilogo	65
6 Discussione e conclusioni	67
6.1 Valore didattico del feedback	67
6.2 Aspetti da considerare nell'applicazione del modulo	67
6.3 Validazione e sviluppi futuri	68
6.4 Considerazioni conclusive	68
Riferimenti bibliografici	70

1 Introduzione

La digitalizzazione della pratica sanitaria ha reso disponibili quantità crescenti e fonti sempre più diversificate di dati clinici. Tuttavia, la loro presenza nei sistemi informativi non determina automaticamente un miglioramento dell'assistenza o del processo decisionale. Per sostenere una decisione clinica, le informazioni rilevanti e le conoscenze mediche devono essere disponibili nel momento in cui il professionista ne ha bisogno.¹

Da questa esigenza deriva lo sviluppo di sistemi capaci di mettere in relazione i dati del paziente con conoscenze cliniche formalizzate, tra cui i Clinical Decision Support Systems e le linee guida cliniche computer-interpretabili. Le Computer-Interpretable Guidelines trasformano il contenuto delle linee guida in rappresentazioni strutturate, espresse mediante modelli e linguaggi elaborabili da un sistema informatico. Tale formalizzazione rende espliciti elementi quali azioni, decisioni e condizioni e consente di utilizzare le raccomandazioni nei sistemi di supporto decisionale.¹⁻³

La presente tesi si colloca nell'ambito di GLARE-Edu, un ambiente formativo basato sull'uso di linee guida cliniche computer-interpretabili. In questo contesto, lo studente può confrontarsi con un pathway clinico e proporre le proprie scelte rispetto al percorso previsto dalla linea guida.⁴ Il lavoro affronta il problema della valutazione di tali risposte, con l'obiettivo di superare un confronto puramente formale tra risposta inserita e risposta attesa.

1.1 Contesto e motivazione: CDSS e linee guida computer-interpretabili

I sistemi di supporto alla decisione clinica, generalmente indicati come *Clinical Decision Support Systems* o CDSS, comprendono strumenti informatici di natura diversa. Alcuni hanno una funzione prevalentemente informativa, altri richiamano l'attenzione dell'utente su possibili criticità attraverso avvisi o allarmi, mentre sistemi più avanzati possono produrre raccomandazioni specifiche sulla base dei dati relativi al paziente. In tutti questi casi, l'obiettivo è assistere il professionista sanitario durante il processo decisionale.¹

L'utilità di questi sistemi è legata anche alla quantità crescente di dati che il professionista deve considerare durante il processo decisionale. Il medico può avere a disposizione molte informazioni, provenienti ad esempio da cartelle cliniche elettroniche, sistemi di prescrizione informatizzata, esami diagnostici e altre fonti cliniche. Tuttavia, durante l'attività quotidiana, non sempre è possibile recuperare e interpretare tutti questi dati in modo tempestivo. Un sistema di supporto alla decisione può essere utile quando aiuta a selezionare le informazioni rilevanti e a metterle in relazione con la conoscenza clinica appropriata.¹

In questo ambito rientrano anche le linee guida cliniche computer-interpretabili, o *Computer-Interpretable Guidelines* (CIGs). Le linee guida tradizionali sono spesso documenti narrativi, pensati per essere letti e interpretati da professionisti sanitari. Per poter essere utilizzate da un sistema informatico, devono invece essere trasformate in una rappresentazione più strutturata, capace di descrivere azioni, condizioni, decisioni e relazioni tra le diverse parti del percorso clinico.^{2,3} La formalizzazione richiede di rendere espliciti gli elementi della linea guida e le relazioni necessarie alla sua elaborazione informatica.

Una linea guida computer-interpretabile non è quindi semplicemente un testo memorizzato in formato digitale, ma una rappresentazione della conoscenza medica che il sistema può elaborare. Le CIG possono costituire la base di sistemi di supporto decisionale.^{2,3}

Nel caso di GLARE-Edu, la stessa rappresentazione strutturata viene impiegata per la navigazione didattica, la simulazione automatica e la verifica delle scelte dello studente. Nella modalità di verifica, le proposte dello studente vengono confrontate passo per passo con le raccomandazioni della linea guida e, in caso di risposta errata, il sistema può fornire una spiegazione. In questo modo, lo studente può esercitarsi non soltanto sul contenuto delle raccomandazioni, ma anche sul modo in cui esse vengono applicate a uno specifico caso clinico.⁴

1.2 Obiettivi e contributo della tesi

La tesi ha come obiettivo la progettazione e l'integrazione in GLARE-Edu di un modulo per la verifica semantica delle risposte dello studente. Il problema nasce dal fatto che, in un ambiente formativo basato su pathway clinici, la valutazione non dovrebbe limitarsi a stabilire se la risposta inserita coincida esattamente con quella prevista dal sistema. In molti casi, infatti, una risposta può essere diversa dal punto di vista formale, ma comunque vicina dal punto di vista del significato clinico.

Un confronto puramente testuale rischia di essere troppo rigido. Nella gerarchia di SNOMED CT, una terminologia clinica strutturata adottata come riferimento nel presente lavoro, due concetti non identici possono essere collegati da relazioni di generalizzazione o specializzazione, oppure condividere antenati comuni nella struttura tassonomica.^{5,6} Per questo motivo, il lavoro propone di considerare la distanza semantica tra i concetti scelti dallo studente e quelli associati agli elementi corretti del pathway.

Il modulo lavora su porzioni delimitate del percorso clinico, in particolare sui segmenti compresi tra una decisione e la successiva. Questa scelta permette di valutare le risposte dello studente rispetto a una parte definita della linea guida, senza considerare l'intero pathway come un unico blocco. Lo studente viene chiamato a individuare gli elementi rilevanti di una determinata porzione del percorso, mentre il sistema confronta tali scelte con quelle previste dal case study.

Per rendere possibile questo confronto, il lavoro utilizza SNOMED CT come riferimento terminologico per rappresentare i concetti clinici coinvolti. SNOMED CT rappresenta i significati clinici mediante concetti, ciascuno dotato di un identificativo numerico univoco, e organizza tali concetti in gerarchie definite da relazioni di tipo *is-a*.⁵ Nel presente lavoro, questa struttura gerarchica viene utilizzata come base per calcolare la distanza semantica. Su questa base, il modulo è progettato per distinguere risposte esatte, concetti più generali o più specifici rispetto a quelli attesi, elementi mancanti ed elementi aggiuntivi.

Il lavoro propone un'estensione di GLARE-Edu orientata a una valutazione più flessibile e informativa delle risposte dello studente. Nell'ambito della valutazione formativa, un feedback che chiarisce gli obiettivi e gli standard attesi, fornisce informazioni sulla prestazione e aiuta a colmare la distanza tra la prestazione corrente e quella desiderata può sostenere lo sviluppo dell'autoregolazione.⁷

In questa prospettiva, la verifica semantica non segnala soltanto che la risposta differisce da quella attesa, ma indica anche la natura della differenza individuata dal modulo.

2 Contesto applicativo: GLARE-Edu e SNOMED CT

Questo capitolo presenta il contesto applicativo del lavoro di tesi. Dopo aver chiarito la differenza tra linee guida cliniche e clinical pathway, viene descritto GLARE-Edu, l'ambiente in cui è stato integrato il modulo di verifica semantica. L'attenzione si concentra poi sulla rappresentazione del percorso clinico e sull'impiego di SNOMED CT per associare i nodi a concetti clinici strutturati, elementi necessari per comprendere il metodo proposto nei capitoli successivi.

2.1 Clinical practice guidelines e pathway clinici

Le linee guida cliniche sono dichiarazioni sviluppate in modo sistematico per assistere professionisti sanitari e pazienti nelle decisioni relative all'assistenza appropriata in specifiche circostanze cliniche. In questo senso, organizzano la conoscenza medica in forma di raccomandazioni utilizzabili nella pratica e forniscono un riferimento generale per orientare le scelte cliniche.⁸

Una linea guida clinica mantiene però, per sua natura, un carattere generale. Essa può indicare quali comportamenti siano raccomandati in una determinata situazione, senza descrivere necessariamente l'intero processo operativo con cui l'assistenza viene organizzata in uno specifico contesto. Nella pratica, l'applicazione delle raccomandazioni deve tenere conto delle condizioni del paziente e dell'organizzazione del percorso di cura.^{8,9}

Per questo motivo è utile distinguere tra linea guida clinica e pathway clinico. Un clinical pathway è un piano assistenziale strutturato e multidisciplinare che traduce raccomandazioni o altre evidenze in un processo di cura applicabile a uno specifico contesto organizzativo, dettagliando i passaggi essenziali del trattamento o dell'episodio assistenziale. Mentre la linea guida fornisce il riferimento conoscitivo generale, il pathway ne organizza l'applicazione operativa per uno specifico problema clinico e una popolazione definita.⁹

Nel contesto della tesi, la linea guida non è considerata soltanto come documento testuale, ma come un percorso clinico strutturato, rappresentabile e utilizzabile all'interno di un ambiente informatico.

2.2 Il sistema GLARE e l'ambiente GLARE-Edu

GLARE, acronimo di *Guideline Acquisition, Representation and Execution*, è un sistema sviluppato per acquisire, rappresentare ed eseguire linee guida cliniche. Una caratteristica rilevante del sistema è la sua indipendenza dal dominio clinico: il formalismo adottato non è progettato per una sola patologia, ma può essere applicato a linee guida appartenenti ad ambiti medici differenti.¹⁰

Il sistema si colloca nell'ambito dell'intelligenza artificiale applicata alla medicina. In particolare, GLARE utilizza tecniche di rappresentazione della conoscenza e di supporto al ragionamento per rendere le linee guida utilizzabili da un sistema informatico. L'obiettivo non è soltanto digitalizzare un documento, ma permettere che la linea guida venga acquisita in forma strutturata, controllata e applicata a un caso specifico.¹⁰

L'organizzazione modulare di GLARE permette di distinguere le fasi di acquisizione, rappresentazione ed esecuzione della linea guida. Il modulo di acquisizione consente agli esperti clinici di introdurre la struttura della linea guida e le proprietà dei suoi elementi attraverso un'interfaccia grafica progettata per essere utilizzabile dagli esperti del dominio. L'ambiente fornisce inoltre funzioni di navigazione e controlli sintattici, logici e temporali sulla coerenza della rappresentazione.¹⁰

Il modulo di esecuzione permette invece di applicare una linea guida acquisita a uno specifico paziente. Durante l'esecuzione, il sistema considera i dati del paziente e lo stato corrente del percorso, interagendo con l'utente attraverso un'interfaccia grafica. In questo modo la linea guida formalizzata può essere utilizzata come supporto operativo all'applicazione del percorso clinico.¹⁰

GLARE-Edu è costruito a partire da GLARE e ne mantiene il formalismo di rappresentazione, spostando l'attenzione dal supporto alla pratica clinica alla formazione. Il sistema è indipendente dal dominio e consente di acquisire in formato computabile sia linee guida sia case study, così da esercitare lo studente sull'applicazione delle migliori pratiche a casi specifici.⁴

Sulla base dei contenuti acquisiti, GLARE-Edu offre tre funzioni principali. La navigazione consente di esplorare la rappresentazione strutturata della linea guida, la simulazione automatica mostra passo per passo come la linea guida suggerisce di gestire un caso, la verifica chiede allo studente di proporre le proprie scelte e le confronta progressivamente con quelle previste, fornendo spiegazioni quando emergono differenze.⁴

GLARE-Edu costituisce l'ambiente applicativo nel quale è stato integrato il modulo di verifica semantica delle risposte dello studente.

2.3 Rappresentazione del pathway: nodi, archi e decisioni

La rappresentazione del pathway in GLARE-Edu deriva dal formalismo adottato in GLARE. Molti approcci alle linee guida computer-interpretabili adottano un *Task-Network Model*, cioè una rappresentazione gerarchica del flusso come rete di attività.² In GLARE, il percorso è rappresentato come una rete gerarchica di azioni: i nodi corrispondono alle azioni della linea guida, mentre le relazioni di controllo stabiliscono quali azioni possono essere eseguite successivamente e in quale ordine.¹⁰

La rappresentazione a grafo consente di modellare sia le sequenze lineari sia i punti di scelta presenti in un pathway clinico. In un percorso diagnostico-terapeutico, infatti, non tutte le azioni seguono semplicemente una dopo l'altra: alcune dipendono dai dati disponibili, altre dal risultato di una decisione, altre ancora possono essere ripetute o soggette a vincoli temporali. Il grafo permette di rendere esplicita la struttura del processo.¹⁰

GLARE distingue innanzitutto tra azioni atomiche e azioni composite. Le azioni atomiche rappresentano passaggi elementari della linea guida, che non richiedono ulteriore decomposizione per essere eseguiti. Le azioni composite, invece, sono definite attraverso le loro componenti e permettono di descrivere parti più ampie della linea guida secondo una relazione gerarchica. In questo modo è possibile rappresentare il percorso a diversi livelli di dettaglio: dalla struttura generale fino ai singoli passaggi operativi.¹⁰

Le azioni atomiche sono ulteriormente distinte in categorie. Le *work actions* rappresentano attività da eseguire in un determinato punto della linea guida e possono essere descritte attraverso attributi come nome, descrizione, costo, tempo, risorse e obiettivi. Le *query actions* sono richieste di informazioni, che possono provenire dal medico, dalla cartella clinica del paziente, da esami di laboratorio o da altre fonti. Le *pharmacological prescriptions* rappresentano le prescrizioni farmacologiche previste dalla linea guida e identificano i farmaci da somministrare. Le *decision actions* rappresentano i criteri utilizzati per selezionare tra percorsi alternativi. Le *conclusions*, infine, costituiscono l'esito esplicito di un processo decisionale.¹⁰

Gli archi del grafo stabiliscono quali azioni possono essere eseguite successivamente e in quale ordine. GLARE prevede diverse relazioni di controllo, tra cui sequenza, relazioni vincolate temporalmente, alternative e ripetizioni. La sequenza indica il normale avanzamento da un'azione alla successiva, l'alternativa rappresenta una diramazione del percorso, la ripetizione consente di modellare attività che possono essere eseguite più volte, i vincoli temporali permettono di descrivere relazioni di durata o intervalli tra azioni.¹⁰

Le decisioni hanno un ruolo particolarmente importante. Esse permettono di scegliere tra rami alternativi del percorso sulla base dei dati disponibili. In GLARE sono previste decisioni diagnostiche e decisioni terapeutiche. Le decisioni diagnostiche considerano parametri relativi allo stato del paziente per discriminare tra ipotesi diverse e possono essere rappresentate anche mediante punteggi e soglie. Le decisioni terapeutiche riguardano invece la scelta tra trattamenti alternativi e sono basate su criteri come efficacia, costo, effetti collaterali, compliance e durata.¹⁰

Nel contesto formativo di GLARE-Edu, questa struttura può essere navigata e utilizzata nella simulazione e nella verifica per mostrare come le diverse azioni e decisioni siano collegate lungo il percorso clinico.⁴ Per la tesi, la rappresentazione a nodi e archi è fondamentale perché consente di individuare porzioni del pathway e di riferire le risposte dello studente a elementi precisi della linea guida, invece di trattarle come risposte isolate.

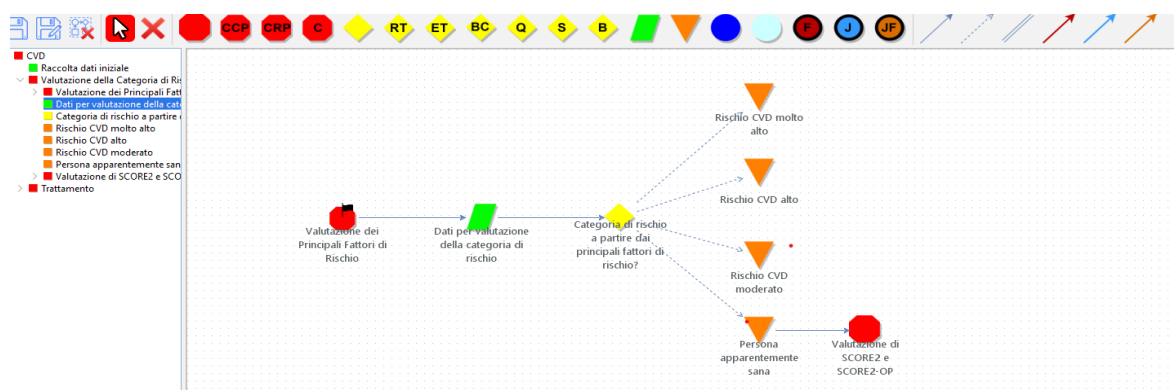


Figura 1. Esempio di rappresentazione grafica di un pathway in GLARE, con azioni, richiesta di dati, decisione e rami alternativi.

2.4 SNOMED CT: concetti, identificativi e gerarchia tassonomica

SNOMED CT è una terminologia clinica strutturata che consente di rappresentare contenuti sanitari in modo coerente ed elaborabile da applicazioni informatiche. Organizza significati clinici e definizioni logiche attraverso concetti, descrizioni e relazioni.^{5,6} Nel contesto del presente lavoro viene utilizzato per associare gli elementi clinici del pathway e le risposte dello studente a riferimenti terminologici identificabili in modo univoco.

L'unità fondamentale di SNOMED CT è il concetto. Un concetto rappresenta un significato clinico e appartiene a una delle principali gerarchie della terminologia, tra cui *clinical finding*, *procedure*, *substance*, *body structure* e *observable entity*. Ogni concetto è associato a un identificativo numerico univoco, chiamato SCTID, che permette di riferirsi al significato clinico indipendentemente dalla forma linguistica usata per descriverlo. Lo stesso significato clinico può essere espresso mediante sinonimi o formulazioni linguistiche differenti.^{5,11}

Accanto ai concetti, SNOMED CT utilizza descrizioni e relazioni. Le descrizioni sono i termini leggibili dall'uomo associati a un concetto, esse permettono di denominare il concetto e di gestire diverse espressioni linguistiche riferite allo stesso significato clinico. Le relazioni, invece, collegano i concetti tra loro e contribuiscono a definirne il significato all'interno della terminologia. In questo modo, SNOMED CT non si presenta come una semplice lista di termini, ma come una struttura organizzata di concetti clinici collegati.^{5,6}

La distinzione tra concetto e descrizione è particolarmente utile nel contesto della verifica semantica. Se si confrontassero soltanto le stringhe testuali, due risposte formulate in modo diverso potrebbero essere considerate differenti anche quando fanno riferimento allo stesso significato clinico. L'uso di concetti identificati in modo univoco permette invece di separare il significato clinico dalle diverse forme linguistiche con cui esso può essere espresso.^{5,6}

Una parte centrale di SNOMED CT è la sua struttura gerarchica. I concetti sono organizzati secondo relazioni di generalizzazione e specializzazione: i concetti più generali si trovano a livelli superiori della gerarchia, mentre quelli più specifici si collocano a livelli inferiori. La relazione principale che costruisce questa organizzazione è la relazione *is-a*, attraverso la quale un concetto viene definito come sottotipo di un concetto più generale.⁶

Questa struttura gerarchica consente di analizzare i rapporti tra concetti clinici. Due concetti possono non coincidere, ma essere comunque collegati: uno può essere più generale dell'altro, più specifico, oppure entrambi possono condividere un antenato comune.⁶ Nel presente lavoro, la posizione dei concetti nella gerarchia costituisce la base per valutare relazioni semantiche che non emergerebbero da un semplice confronto testuale.

SNOMED CT ha inoltre una struttura poli-gerarchica: uno stesso concetto può avere più concetti sovraordinati e appartenere a più percorsi gerarchici.⁶ Questa caratteristica richiede attenzione nell'analisi di antenati, discendenti e concetti comuni, perché tra un concetto specifico e i livelli più generali della terminologia possono esistere più cammini.

Oltre alle relazioni gerarchiche di tipo *is-a*, SNOMED CT comprende attributi, o tipi di relazione, impiegati per rappresentare caratteristiche del significato di un concetto. Il modello concettuale stabilisce inoltre i domini nei quali ciascun attributo può essere applicato e i valori ammessi.¹¹

Il modulo sviluppato utilizza soprattutto la struttura gerarchica di SNOMED CT, sulla quale si basa il confronto tra la risposta dello studente e il concetto previsto dal pathway.

Nel presente lavoro, SNOMED CT svolge una duplice funzione: da un lato, consente di rappresentare i concetti clinici mediante identificativi standardizzati e univoci, dall'altro, fornisce una struttura gerarchica attraverso cui analizzare le relazioni semantiche tra i concetti stessi. Tali caratteristiche costituiscono la base terminologica per i metodi di confronto semantico descritti nel capitolo successivo.

2.5 Collegamento tra nodi GLARE e concetti SNOMED

Nel pathway rappresentato in GLARE-Edu, gli elementi della linea guida hanno innanzitutto una funzione strutturale. Essi descrivono azioni atomiche o composite, richieste di informazioni, decisioni o conclusioni e sono collegati da relazioni che governano il flusso.¹⁰ Questa rappresentazione è necessaria per descrivere la logica della linea guida, ma non è sufficiente, da sola, per valutare il significato clinico delle risposte inserite dallo studente.

La verifica semantica richiede che gli elementi clinici del pathway siano associati a concetti espliciti. A questo scopo viene utilizzato SNOMED CT: ciascun nodo verificabile è descritto non soltanto dall'etichetta testuale e dalla posizione nel percorso, ma anche da un concetto clinico identificato in modo univoco.

Il collegamento tra nodi GLARE e concetti SNOMED permette di distinguere due livelli della rappresentazione. Il primo livello è strutturale e riguarda il ruolo del nodo nel pathway: la sua posizione, il tipo di elemento rappresentato e le relazioni di controllo con gli altri nodi.¹⁰ Il secondo livello è semantico e riguarda il significato clinico associato a quel nodo e la sua collocazione nella gerarchia SNOMED CT.^{5,6} La verifica sviluppata nella tesi utilizza entrambi i livelli: il pathway permette di individuare quali elementi devono essere valutati in un determinato segmento, mentre SNOMED CT permette di confrontare il significato clinico di tali elementi con le risposte dello studente.

Non tutti i nodi del pathway sono trattati nello stesso modo. Nel modulo sviluppato, gli elementi considerati semanticamente verificabili sono soprattutto quelli che corrispondono a contenuti clinici direttamente confrontabili, come azioni, richieste di dati e prescrizioni farmacologiche, quando tali elementi sono associati a concetti SNOMED CT. Le decisioni e le conclusioni hanno invece una funzione diversa: orientano il percorso o rappresentano l'esito di una scelta, e vengono trattate separatamente rispetto al confronto semantico dei singoli concetti clinici.

L'integrazione tra il pathway di GLARE-Edu e la terminologia SNOMED CT consente di riferire le risposte dello studente a elementi precisi della linea guida e costituisce la base dei metodi di verifica semantica descritti nel capitolo successivo.

3 Metodi per la verifica semantica

Nei capitoli precedenti sono stati introdotti il contesto applicativo di GLARE-Edu e il ruolo di SNOMED CT come terminologia clinica di riferimento. A partire da tali premesse, il problema affrontato nel presente lavoro può essere formulato in termini più precisi: come confrontare le risposte dello studente con gli elementi previsti dal pathway quando le rispettive formulazioni non coincidono necessariamente sul piano testuale.

Un confronto basato esclusivamente sull'uguaglianza tra stringhe risulterebbe troppo limitato. In ambito clinico, infatti, concetti differenti possono essere collegati da relazioni di generalità, specificità o vicinanza semantica. La valutazione richiede pertanto metodi in grado di considerare non soltanto l'etichetta associata a ciascun concetto, ma anche la sua posizione all'interno della struttura terminologica.

Il capitolo presenta i principali riferimenti teorici impiegati per sviluppare il metodo di verifica semantica. In primo luogo, vengono introdotte le nozioni di similarità e distanza semantica. Successivamente, viene esaminato il ruolo dell'Information Content e dell'antenato comune più informativo, elementi centrali nelle misure basate su tassonomie. Le sezioni successive applicano tali nozioni alla misura adottata e al problema dell'abbinamento tra le risposte fornite e le soluzioni attese.

3.1 Similarità e distanza semantica tra concetti

Le misure semantiche nascono dall'esigenza di confrontare entità sulla base del loro significato. A differenza delle misure puramente sintattiche, che valutano la somiglianza tra forme testuali, le misure semantiche cercano di stimare il grado di relazione tra parole, concetti o istanze considerando informazioni che ne descrivono il significato. Questa distinzione è importante perché la forma linguistica non coincide sempre con il contenuto concettuale: due espressioni diverse possono rimandare a concetti vicini, mentre espressioni simili possono indicare significati differenti.¹²

La letteratura distingue diverse nozioni collegate tra loro. La *relatedness* semantica indica una relazione ampia tra due elementi: due concetti possono essere collegati perché appartengono allo stesso dominio, perché compaiono nello stesso contesto o perché uno richiama l'altro. La similarità semantica è invece più ristretta e riguarda soprattutto il grado in cui due concetti condividono caratteristiche comuni. Quando il confronto avviene all'interno di una tassonomia, la similarità è spesso legata alle relazioni di tipo gerarchico, cioè ai rapporti tra concetti più generali e concetti più specifici. La stessa relazione di vicinanza può essere espressa anche mediante la distanza semantica: più due concetti sono semanticamente vicini, minore sarà la distanza tra essi.^{12,13}

Nel presente lavoro, l'interesse principale riguarda misure basate sulla conoscenza rappresentata in una terminologia strutturata. In questo tipo di approccio, il confronto utilizza esplicitamente i concetti, le loro proprietà e le relazioni rappresentate in una tassonomia o in un'ontologia.¹² Poiché il modulo utilizza SNOMED CT, la valutazione delle risposte rientra in questa impostazione *knowledge-based*.

Un modo immediato per confrontare due concetti in una gerarchia consiste nel contare il numero di archi che li separa. Se due concetti sono collegati da un percorso breve, vengono considerati più vicini, se il percorso è lungo, vengono considerati più distanti. Questo approccio è intuitivo, ma presenta un limite significativo: assume che tutti gli archi della gerarchia abbiano lo stesso peso semantico. Nelle tassonomie reali questa assunzione è difficile da mantenere, perché alcune parti della struttura possono essere molto dettagliate, mentre altre possono essere più generali.^{12,13}

Resnik evidenzia proprio questo problema discutendo le misure basate sul conteggio degli archi. In una tassonomia, due collegamenti collocati in zone diverse non rappresentano necessariamente la stessa differenza concettuale. Un arco può separare due concetti molto vicini in una parte densa della gerarchia, mentre in un'altra zona può collegare concetti tra loro più distanti. Per questo motivo, la semplice lunghezza del percorso non sempre restituisce una misura adeguata della vicinanza semantica.¹³

Una considerazione simile emerge anche nel contesto delle ontologie biomedicali. Pesquita et al. distinguono gli approcci edge-based, che utilizzano principalmente le relazioni del grafo, dagli approcci node-based, che considerano proprietà dei nodi, dei loro antenati o delle annotazioni associate. Gli approcci edge-based hanno il vantaggio di essere semplici, ma possono risentire della distribuzione non uniforme dei nodi e degli archi. Gli approcci node-based cercano invece di rappresentare la specificità dei concetti in modo meno dipendente dalla sola forma della gerarchia.¹²

Tra gli approcci node-based rientrano le misure basate sull'Information Content. L'idea di fondo è che un concetto molto generale sia poco informativo, perché comprende molti casi diversi, mentre un concetto più specifico sia più informativo, perché identifica una porzione più ristretta del dominio. In questo modo, la valutazione della vicinanza tra concetti non dipende soltanto dal numero di passaggi nella gerarchia, ma anche dal contenuto informativo associato ai concetti coinvolti.^{12,13}

Questa impostazione è adatta al problema considerato, poiché consente di quantificare la relazione semantica tra il concetto selezionato dallo studente e quello atteso. Nel modulo sviluppato, la distanza è impiegata come costo nella costruzione della matrice delle distanze, utilizzata per individuare l'abbinamento complessivamente migliore tra risposte e soluzioni. La classificazione didattica delle coppie selezionate viene applicata in una fase successiva, secondo le regole descritte nel Capitolo 4.

3.2 Information Content e Lowest Common Subsumer

L'Information Content (IC) esprime quanto un concetto sia informativo all'interno di una tassonomia. Resnik associa a ogni concetto una probabilità $p(c)$: un concetto generale comprende anche le istanze dei propri sottotipi e tende ad avere probabilità più alta e IC più basso, un concetto specifico ha invece probabilità minore e IC più elevato.¹³

L'Information Content di un concetto viene espresso dalla formula:

$$IC(c) = -\log p(c)^{13}$$

La relazione tra probabilità e informazione è inversa: in presenza di una radice unica, il concetto più generale ha probabilità pari a 1 e IC pari a 0. Rispetto alla sola profondità o al conteggio degli archi, l'IC consente di rappresentare meglio la specificità anche in zone della tassonomia con densità diversa.^{12,13}

Per confrontare due concetti è necessario considerare non soltanto il rispettivo Information Content, ma anche l'informazione condivisa, rappresentata dai loro antenati comuni. Un antenato comune con un valore di IC elevato indica una maggiore vicinanza semantica.¹³

Il Lowest Common Subsumer (LCS) è l'antenato comune più specifico dei due concetti confrontati. In una gerarchia ad albero è univoco, mentre nelle strutture con ereditarietà multipla possono esistere percorsi differenti e più antenati comuni candidati.^{12,13}

Questo aspetto è rilevante nel caso di SNOMED CT, che presenta una struttura poli-gerarchica: uno stesso concetto può avere più genitori e due concetti possono essere collegati attraverso diversi percorsi gerarchici. L'individuazione degli antenati comuni deve pertanto tenere conto dei diversi percorsi presenti nella gerarchia.⁶

La misura proposta da Resnik rappresenta la similarità tra due concetti mediante l'Information Content del loro antenato comune più informativo, indicato come Most Informative Common Ancestor (MICA):

$$sim_{Resnik}(c1, c2) = IC(MICA(c1, c2))$$

Due concetti risultano più simili quando condividono un antenato con un valore di IC elevato. La misura di Resnik considera però soltanto l'informazione condivisa attraverso tale antenato e non tiene conto separatamente dell'Information Content dei due concetti confrontati.^{12,13}

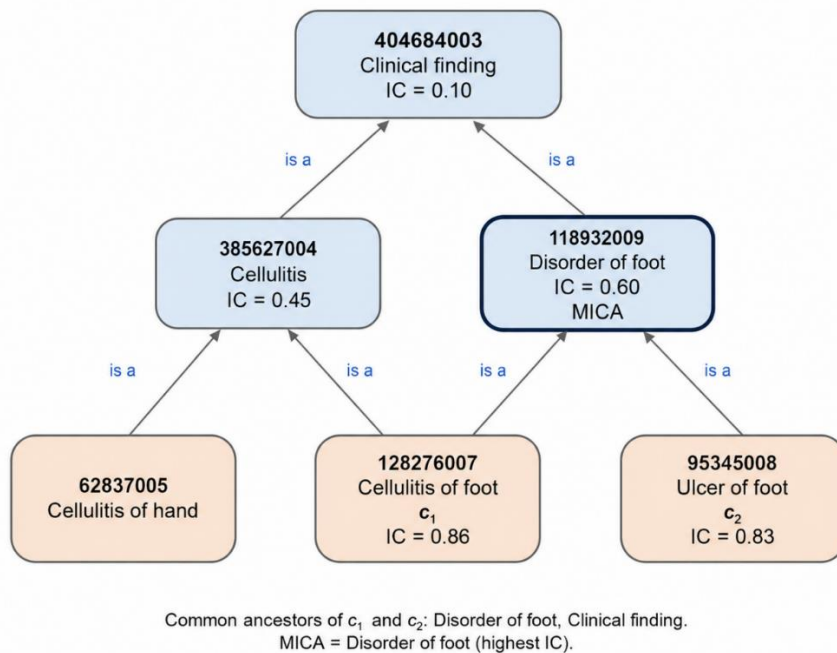


Figura 2. Esempio di individuazione del Most Informative Common Ancestor (MICA) in una struttura tassonomica, con valori di Information Content illustrativi.

Nelle ontologie biomedicali, l'IC può essere stimato a partire dalle frequenze di annotazione oppure da caratteristiche interne della struttura ontologica. Le misure basate sull'IC sono meno sensibili ai problemi derivanti dalla densità variabile del grafo, tuttavia, le stime fondate sulle annotazioni possono essere influenzate anche dalla distribuzione e dalla disponibilità dei dati.¹²

La misura di Resnik mostra come l'Information Content dell'antenato comune possa rappresentare l'informazione condivisa tra due concetti. La distanza Jiang-Conrath, descritta nella sezione successiva, considera anche l'Information Content dei due concetti confrontati e calcola la distanza utilizzando l'IC del loro Lowest Common Subsumer.

3.3 Distanza Jiang-Conrath

La distanza di Jiang-Conrath nasce dal tentativo di combinare due modi diversi di misurare la vicinanza tra concetti in una tassonomia. Da un lato vi sono gli approcci basati sull'Information Content, che attribuiscono importanza alla specificità dei concetti e alla quantità di informazione condivisa. Dall'altro vi sono gli approcci basati sugli archi, che valutano la distanza considerando il percorso che collega due nodi nella struttura gerarchica.¹⁴

Jiang e Conrath osservano che entrambi gli approcci presentano vantaggi e limiti. Le misure basate solo sul conteggio degli archi sono semplici e intuitive, ma presuppongono che ogni arco della tassonomia rappresenti la stessa distanza semantica. Questa ipotesi è problematica nelle tassonomie reali, dove alcune aree possono essere molto dense e dettagliate, mentre altre possono essere più generali. Al contrario, l'Information Content permette di rappresentare meglio la specificità dei concetti, ma da solo non descrive sempre in modo completo la distanza tra i concetti confrontati.¹⁴

L'idea centrale proposta dagli autori è interpretare la distanza tra un concetto figlio e il suo concetto padre come differenza tra i rispettivi valori di Information Content. Se un concetto figlio è molto più specifico del padre, la differenza informativa tra i due sarà maggiore, se invece il passaggio dal padre al figlio introduce poca informazione aggiuntiva, la distanza sarà minore. In questo modo, il peso di un collegamento nella tassonomia non viene considerato uniforme, ma dipende dall'informazione che distingue il concetto più specifico da quello più generale.¹⁴

A partire da questa intuizione, la distanza tra due concetti viene calcolata considerando il percorso che li collega attraverso il Lowest Common Subsumer (LCS). Nella componente basata sull'Information Content, la distanza Jiang-Conrath può essere espressa come:

$$\text{dist}_{\{IC\}}(c_1, c_2) = IC(c_1) + IC(c_2) - 2 * IC(LCS(c_1, c_2))$$

Dove c_1 e c_2 sono i concetti confrontati, mentre $LCS(c_1, c_2)$ indica il loro antenato comune rilevante nella tassonomia. La formula misura quanta informazione appartiene ai due concetti ma non è condivisa attraverso l'antenato comune. Se due concetti sono vicini, l'antenato comune conserva una parte significativa della loro informazione e la distanza risulta bassa. Se invece i concetti condividono solo un antenato molto generale, l'informazione comune è minore e la distanza aumenta.¹⁴

Rispetto alla misura di Resnik, che considera soltanto l'Information Content dell'antenato comune più informativo, la distanza Jiang-Conrath tiene conto anche dei due concetti di partenza. Questo aspetto è importante perché due coppie di concetti possono condividere lo stesso antenato comune, ma trovarsi a distanze diverse da esso. Considerare solo l'informazione condivisa non basta sempre a distinguere questi casi, la formula Jiang-Conrath considera anche la parte di informazione propria dei due concetti che non viene coperta dall'antenato comune.^{13,14}

La misura può essere letta come una distanza informativa: non conta soltanto quanto due concetti condividono, ma anche quanto ciascuno di essi si allontana dalla parte comune. Quando i concetti coincidono, la distanza tende al valore minimo, perché l'informazione dei concetti e quella dell'antenato comune coincidono. Quando invece i concetti appartengono a rami diversi della gerarchia, l'antenato comune tende a essere più generale e il valore della distanza cresce.¹⁴

Nell'articolo, Jiang e Conrath valutano la misura utilizzando WordNet come tassonomia lessicale e frequenze ricavate da corpora per stimare l'Information Content. Il loro obiettivo è mostrare che una misura capace di combinare struttura tassonomica e informazione statistica può rappresentare meglio la distanza semantica rispetto a un semplice conteggio degli archi. Questa impostazione è rilevante anche per il presente lavoro, perché la verifica semantica richiede una misura in grado di distinguere concetti identici, concetti vicini e concetti semanticamente più distanti.¹⁴

La distanza Jiang-Conrath fornisce una base metodologica per misurare la separazione informativa tra due concetti all'interno di una tassonomia. Il suo punto centrale è considerare sia l'Information Content dei concetti confrontati sia l'informazione condivisa attraverso un antenato comune. Il modo in cui l'Information Content viene calcolato nel contesto della gerarchia SNOMED CT viene descritto nella sezione successiva.¹⁴

3.4 Adattamento dell'Information Content alla gerarchia SNOMED CT

Per utilizzare la distanza semantica descritta nella sezione precedente è necessario stabilire come calcolare l'Information Content dei concetti. Nelle misure classiche basate su IC, il valore informativo di un concetto è definito a partire dalla sua probabilità di occorrenza: un concetto molto frequente porta poca informazione, mentre un concetto raro o più specifico risulta più informativo.¹³⁻¹⁵

La misura proposta da Seco, Veale e Hayes segue un'impostazione diversa. Invece di ricavare l'Information Content da frequenze osservate in raccolte di testi, utilizza direttamente la struttura della tassonomia. Il presupposto è che la gerarchia stessa contenga informazioni sulla generalità o specificità dei concetti. Un concetto che ha molti iponimi è più generale, un concetto foglia, che non viene ulteriormente specializzato, è invece più specifico.¹⁵

Nel lavoro originale questa idea viene applicata a WordNet. Il valore di Information Content è calcolato in funzione del numero di iponimi del concetto, indicati nella formula come $hypo(c)$:¹⁵

$$ic_{wn}(c) = \frac{\log(\frac{hypo(c)+1}{max_{wn}})}{\log(\frac{1}{max_{wn}})} = 1 - \frac{\log(hypo(c) + 1)}{\log(max_{wn})}$$

dove $hypo(c)$ indica il numero di iponimi del concetto c , mentre max_{wn} indica il numero massimo di concetti presenti nella tassonomia considerata. La formula assegna valori più bassi ai concetti con molti discendenti e valori più alti a quelli con pochi o nessun discendente. Se un concetto non presenta discendenti, il numeratore assume valore nullo e l'Information Content raggiunge il valore massimo.¹⁵

Nel presente lavoro, il principio della misura di Seco, Veale e Hayes viene adattato alla gerarchia SNOMED CT: gli iponimi della formulazione originaria vengono sostituiti dai discendenti distinti individuati attraverso le relazioni *is-a*. In questo modo, il numero di discendenti diventa l'indicatore usato per stimare il grado di generalità o specificità di ciascun concetto nella gerarchia SNOMED CT.

La formula utilizzata diventa:

$$IC = 1 - \frac{\log(desc(c) + 1)}{\log(N)}$$

dove $desc(c)$ indica il numero di discendenti del concetto c , mentre N indica il numero complessivo di concetti considerati nella gerarchia. Il comportamento della formula rimane lo stesso: un concetto con molti discendenti è considerato più generale e riceve un valore di IC più basso, un concetto con pochi discendenti, o senza discendenti, è considerato più specifico e riceve un valore più alto.

Poiché la gerarchia di SNOMED CT è poli-gerarchica e un concetto può avere più genitori, uno stesso discendente può essere raggiunto attraverso più cammini gerarchici.⁶ Per questo motivo, nel calcolo implementato vengono considerati soltanto i discendenti distinti, evitando che lo stesso concetto sia contato più volte. Questa scelta è specifica del presente lavoro e consente di applicare la misura strutturale senza introdurre duplicazioni dovute ai diversi cammini gerarchici.

I valori di Information Content così ottenuti vengono poi usati nel calcolo della distanza semantica tra i concetti.

3.5 Problema di assegnamento e Hungarian Algorithm

Dopo aver definito una misura per calcolare la distanza tra concetti resta un secondo problema metodologico: scegliere le corrispondenze migliori quando sono disponibili più possibili abbinamenti. Questo problema può essere ricondotto al problema di assegnamento, nel quale occorre associare gli elementi di un insieme agli elementi di un altro insieme rispettando un vincolo di corrispondenza uno-a-uno.^{16,17}

Il problema può essere risolto mediante l'algoritmo ungherese, noto anche come algoritmo di Kuhn-Munkres, che individua un assegnamento ottimale rispetto ai valori contenuti nella matrice.

Nella formulazione classica, il problema considera n individui e n lavori. Per ogni possibile coppia individuo-lavoro è disponibile un valore numerico, detto rating, che indica la qualità o il valore di quella specifica assegnazione. Tali valori sono raccolti in una matrice:^{16,17}

$$R = (r_{ij})$$

dove r_{ij} rappresenta il valore dell'assegnamento dell'individuo i al lavoro j . Un assegnamento completo consiste nello scegliere un solo elemento per ogni riga e per ogni colonna della matrice. In altre parole, ogni individuo deve ricevere un lavoro e nessun lavoro può essere assegnato a più individui.¹⁷

L'obiettivo è trovare la combinazione di assegnamenti che rende massima la somma dei valori scelti. Se per ogni riga i viene scelta una colonna j_i , la somma dei rating selezionati è:

$$r_{1j_1} + r_{2j_2} + \dots + r_{nj_n} \quad 16,17$$

Il problema consiste nel trovare l'assegnamento che massimizza questa somma, con il vincolo che le colonne j_i siano tutte diverse tra loro. Questo vincolo è essenziale: non basta scegliere il valore migliore in ogni riga, perché la stessa colonna potrebbe essere scelta più volte. La soluzione deve essere globale e deve tenere conto contemporaneamente di tutte le righe e di tutte le colonne.^{16,17}

Per trasformare il problema in una forma di costo minimo, si considera il valore massimo presente nella matrice dei rating:¹⁷

$$r = \max_{i,j} r_{ij}$$

A partire da questo valore si costruisce una nuova matrice $A = x_{ij}$, definita da:

$$x_{ij} = r - r_{ij}$$

In questo modo, un rating alto nella matrice originaria corrisponde a un valore basso nella matrice trasformata. Massimizzare la somma dei rating nella matrice R diventa quindi equivalente a minimizzare la somma degli elementi scelti nella matrice A . Il problema può allora essere espresso come scelta di n elementi indipendenti, cioè elementi che non appartengono alla stessa riga né alla stessa colonna, con somma minima.¹⁷

L'algoritmo lavora sulla matrice di costo trasformandola progressivamente senza modificare il problema di assegnamento sottostante. In una prima fase vengono sottratti i valori minimi delle righe e delle colonne, in modo da ottenere una matrice con valori non negativi e con zeri in posizioni utili per costruire l'assegnamento. Gli zeri rappresentano infatti possibili scelte a costo minimo nella matrice trasformata.¹⁷

A questo punto si cerca un insieme di zeri indipendenti. Se è possibile scegliere n zeri tali che nessuno condivida la stessa riga o la stessa colonna, questi zeri individuano un assegnamento ottimo. Se invece gli zeri disponibili non permettono ancora di costruire un assegnamento completo, la matrice viene modificata ulteriormente.¹⁷

La modifica avviene considerando le linee che coprono gli zeri presenti nella matrice. Quando il numero di linee non è sufficiente a ottenere un assegnamento completo, si individua il più piccolo elemento non coperto e si usa questo valore per trasformare nuovamente la matrice. L'effetto della trasformazione è creare nuovi zeri, mantenendo però equivalente il problema di assegnamento. Il procedimento continua finché non è possibile individuare un insieme completo di zeri indipendenti.¹⁷

Il motivo per cui è necessario un algoritmo di questo tipo è la crescita molto rapida del numero di assegnamenti possibili. In una matrice $n \times n$, gli assegnamenti completi sono $n!$. Valutarli uno per uno sarebbe impraticabile già per valori non molto grandi di n . L'Hungarian Algorithm permette invece di arrivare a una soluzione ottima attraverso trasformazioni successive della matrice.¹⁷

Nel contesto della tesi, la matrice utilizzata ha già il significato di matrice di costi, perché i suoi valori rappresentano distanze semantiche tra concetti. Per questo motivo non è necessario partire da rating da massimizzare: il problema è direttamente quello di scegliere una combinazione di corrispondenze che renda minima la somma complessiva delle distanze, rispettando lo stesso vincolo di assegnamento uno-a-uno. L'applicazione concreta al confronto tra risposte dello studente e soluzioni attese sarà descritta nel capitolo successivo.

4 Progettazione e sviluppo del modulo

Il capitolo descrive la progettazione e lo sviluppo del modulo di verifica semantica integrato in GLARE-Edu. Dopo l'introduzione dei riferimenti teorici nel Capitolo 3, l'attenzione si sposta sull'applicazione concreta di tali strumenti nel flusso di esercitazione: individuazione del segmento, acquisizione delle risposte, estrazione della soluzione corretta, calcolo della distanza, matching, classificazione, score e feedback.

Il modulo non sostituisce il funzionamento generale di GLARE-Edu. Interviene sugli elementi del pathway che possono essere collegati a concetti SNOMED CT e li valuta rispetto alla soluzione prevista dal case study. La verifica viene svolta per porzioni del percorso, mantenendo il legame con la struttura della guideline e con la fase del ragionamento clinico affrontata dallo studente.

Nel capitolo vengono distinti due casi. Le work action sono trattate come elementi a concetto singolo, mentre data enquiry e pharmacological prescription sono trattate come elementi composti da parametri SNOMED CT. Questa distinzione influenza sia la modalità di inserimento delle risposte sia il modo in cui vengono interpretati i risultati della valutazione.

4.1 Dal flusso originale alla verifica semantica per segmenti

Nel flusso originario di GLARE-Edu, la verifica dell'apprendimento è collegata all'avanzamento dello studente lungo il pathway clinico e alle risposte richieste dal case study. Lo studente propone progressivamente come gestire il caso e il sistema confronta le sue scelte, passo per passo, con le raccomandazioni della linea guida.⁴

Il modulo sviluppato introduce una modalità di verifica dedicata agli elementi del pathway che possono essere rappresentati mediante concetti SNOMED CT. Il percorso corretto non viene trattato come un unico blocco, ma viene suddiviso in segmenti delimitati dai punti decisionali della guideline. Ogni segmento corrisponde a una porzione circoscritta del percorso clinico su cui lo studente deve concentrarsi.

Per ciascun segmento vengono considerati gli elementi semanticamente verificabili: work action, data enquiry e pharmacological prescription. Le decisioni mantengono il loro ruolo nel controllo del flusso e nella delimitazione del percorso, ma non sono oggetto diretto del confronto semantico.

Lo studente inserisce le risposte relative al segmento corrente selezionando concetti dall'ontologia. Il sistema confronta tali risposte con gli elementi corretti presenti nello stesso segmento, calcola le distanze semantiche e individua gli abbinamenti più appropriati tra risposte e soluzioni. Il risultato viene trasformato in categorie di valutazione e in uno score didattico.

4.2 Estrazione delle soluzioni corrette dal segmento

Una volta individuato un segmento, il modulo ricava gli elementi corretti da usare nel confronto con le risposte dello studente. La soluzione attesa viene estratta dai nodi del segmento appartenenti al pathway corretto del case study, mantenendo la verifica collegata alla guideline già rappresentata in GLARE-Edu.

L'estrazione riguarda solo i tipi di nodo supportati dalla verifica semantica. Le work action vengono trattate come elementi semplici: ogni nodo corretto fornisce un path SNOMED da confrontare con le risposte inserite dallo studente. Il risultato è una lista di path, ciascuno corrispondente a una work action attesa nel segmento.

Le data enquiry richiedono una selezione più specifica. Il modulo considera solo le richieste di dati collegate esplicitamente a SNOMED CT, cioè quelle in cui il riferimento ontologico è presente e utilizzabile per il confronto. I path recuperati non vengono uniti in una lista indistinta, ma rimangono nel gruppo del nodo enquiry a cui appartengono.

Le pharmacological prescription vengono mantenute come gruppi separati. Nel progetto corrente, il riferimento corretto associato a una prescrizione viene recuperato dal nodo come path SNOMED. La valutazione utilizza comunque la logica dei gruppi, così da non mescolare prescrizioni diverse durante il matching. Questa impostazione rende coerente il trattamento delle prescrizioni con quello delle data enquiry, pur rispettando la diversa struttura dei riferimenti recuperati.

Al termine dell'estrazione, il modulo dispone di due forme di soluzione corretta: una lista semplice per le work action e una raccolta di gruppi per data enquiry e pharmacological prescription.

4.3 Acquisizione delle risposte tramite selezione ontologica

Per rendere le risposte confrontabili dal punto di vista semantico, il modulo non utilizza testo libero. Una formulazione libera potrebbe contenere sinonimi o formulazioni differenti dello stesso significato clinico. SNOMED CT distingue i concetti dalle descrizioni e consente di associare più termini allo stesso concetto.^{5,6} L'interazione avviene attraverso una selezione ontologica basata su SNOMED CT.

Lo studente naviga la gerarchia dei concetti e sceglie il concetto ritenuto corretto per il segmento. Ogni selezione viene salvata come path SNOMED, cioè come sequenza di identificativi che descrive il percorso seguito nella gerarchia fino al concetto selezionato. Il path completo viene mantenuto per la visualizzazione del percorso e per il confronto successivo nell'interfaccia.

La schermata di risposta mantiene separate le categorie di elementi valutabili. Per le work action, lo studente aggiunge direttamente uno o più concetti. Per data enquiry e pharmacological prescription, invece, crea prima un elemento strutturato e aggiunge poi i relativi parametri SNOMED CT. L'organizzazione delle risposte riproduce la distinzione già presente nella soluzione corretta: lista semplice per le work action, gruppi di parametri per gli elementi composti.

Ai fini del calcolo semantico viene utilizzato il conceptId finale del path, perché rappresenta il concetto effettivamente scelto. Il path completo resta utile per il feedback, poiché consente di mostrare allo studente il percorso seguito nella gerarchia e l'eventuale divergenza rispetto al path corretto.

Al termine dell'acquisizione, le risposte dello studente hanno la stessa forma delle soluzioni corrette: path singoli per le work action e gruppi di path per data enquiry e pharmacological prescription. Questa uniformità permette di applicare le successive fasi di confronto in modo coerente.

4.4 Calcolo della distanza semantica

Il calcolo della distanza semantica applica al modulo i metodi introdotti nel Capitolo 3. Ogni confronto tra una risposta dello studente e un elemento corretto viene ricondotto al confronto tra due concetti SNOMED CT. Il risultato è un valore numerico che esprime la separazione semantica tra i concetti coinvolti.

Il modulo usa questo valore come base per il matching e per la successiva classificazione. Una distanza pari a zero indica coincidenza tra i concetti confrontati, valori maggiori indicano uno scostamento nella gerarchia SNOMED CT. Il calcolo è svolto sui `conceptId` finali dei path, mentre i path completi sono conservati per la visualizzazione.

4.4.1 Utilizzo della transitive closure

Il calcolo della distanza richiede operazioni frequenti sulla gerarchia SNOMED CT. Per ogni coppia di concetti occorre recuperare gli antenati comuni e conoscere il numero di discendenti utilizzato nel calcolo dell'Information Content. La selezione dell'antenato comune impiegato nella formula di Jiang-Conrath è descritta nella Sezione 4.4.2. Ripercorrere la gerarchia a ogni confronto renderebbe la valutazione meno adatta a un uso interattivo.

Per ridurre questo costo, il modulo utilizza una transitive closure della tassonomia SNOMED CT. La documentazione ufficiale definisce la transitive closure come l'insieme completo delle relazioni tra ogni concetto e tutti i suoi supertipi, comprendendo sia i genitori diretti sia gli antenati. Una tabella di closure consente inoltre di verificare in modo efficiente le relazioni gerarchiche di generalizzazione e specializzazione tra concetti.¹⁸ Nel modulo, questa struttura viene utilizzata per recuperare rapidamente gli antenati di un concetto e, considerando le coppie in direzione inversa, i suoi discendenti distinti.

Le informazioni della transitive closure vengono caricate in memoria una sola volta, prima della valutazione. Questa organizzazione sposta il costo della navigazione gerarchica in una fase preliminare e permette di eseguire i confronti semantici senza rileggere ripetutamente la tassonomia SNOMED CT.

Nel progetto, la transitive closure viene caricata dal file di testo, organizzata come coppie *subtypeld-supertypeId*. Durante il caricamento vengono costruite in memoria le associazioni tra sottotipi e supertipi e i conteggi dei discendenti distinti, che vengono poi riutilizzati nei confronti successivi.

Il valore N usato nella formula dell'Information Content è ricavato dal database SNOMED CT contando i `conceptId` distinti con `active = 1`. Nella configurazione utilizzata il valore è $N = 372\,406$, la normalizzazione considera i concetti attivi presenti nel database.

4.4.2 Calcolo della distanza Jiang-Conrath

Per ogni coppia di concetti, il modulo utilizza la transitive closure per recuperare gli antenati comuni e individuare il Lowest Common Subsumer (LCS) impiegato nella formula di Jiang-Conrath. Poiché SNOMED CT presenta una struttura poli-gerarchica, possono esistere più antenati comuni candidati, tra questi viene selezionato quello con il valore di Information Content più elevato.

Nel seguito, il concetto così individuato viene indicato come LCS, in accordo con la formulazione della distanza Jiang-Conrath, mentre il criterio di selezione corrisponde operativamente alla nozione di Most Informative Common Ancestor (MICA) descritta nella Sezione 3.2.^{6,14}

Come descritto nella Sezione 3.4, l'Information Content viene calcolato adattando la misura intrinseca di Seco, Veale e Hayes alla gerarchia SNOMED CT. Il numero degli iponimi della formulazione originaria viene sostituito dal numero di discendenti distinti del concetto:

$$IC = 1 - \frac{\log(desc(c) + 1)}{\log(N)}$$

dove $desc(c)$ indica il numero di discendenti distinti del concetto c , mentre N è il numero complessivo dei concetti attivi indicato nella Sezione 4.4.1. L'uso dei discendenti distinti evita che uno stesso concetto venga contato più volte quando può essere raggiunto attraverso percorsi gerarchici differenti. Il valore di IC ottenuto è compreso nell'intervallo $([0,1])$.¹⁵

Dopo aver calcolato l'Information Content dei due concetti e del loro LCS, il modulo applica la distanza Jiang-Conrath:

$$dist(c1, c2) = IC(c1) + IC(c2) - 2 * IC(LCS(c1, c2))$$

La distanza è pari a zero quando i concetti coincidono e aumenta al diminuire dell'informazione condivisa attraverso il loro LCS. Poiché i valori di IC sono normalizzati nell'intervallo $([0,1])$, le distanze valide sono comprese tra 0 e 2. Se il confronto non può essere calcolato, ad esempio a causa di un path non valido o dell'assenza di un antenato comune utilizzabile, il modulo assegna il valore convenzionale 3,0. Durante la verifica di un segmento, il calcolo viene eseguito per tutte le possibili coppie formate da una risposta dello studente e da un elemento corretto. I valori ottenuti costituiscono la matrice delle distanze utilizzata nella successiva fase di matching.

4.5 Matching e valutazione delle risposte

A partire dalla matrice delle distanze, il modulo decide quali risposte associare agli elementi corretti del segmento. L'ordine di inserimento non viene usato come criterio di corrispondenza: una risposta può essere abbinata a qualunque elemento corretto, se tale associazione contribuisce alla distanza complessiva minima.

Il problema viene trattato come un problema di assegnamento, secondo l'impostazione discussa nel Capitolo 3. Le distanze Jiang-Conrath costituiscono i costi da minimizzare e il metodo ungherese individua un assegnamento globale a costo complessivo minimo.^{16,17} Il matching definisce le coppie risposta-soluzione, mentre la categoria finale viene assegnata in un passaggio successivo.

Durante l'esecuzione, il metodo ungherese trasforma la matrice di costo mediante operazioni che preservano il problema di assegnamento e utilizza gli zeri della matrice trasformata per costruire una soluzione ottima.¹⁷ Gli zeri così prodotti sono valori operativi dell'algoritmo e non indicano necessariamente una distanza semantica pari a zero.

Per ogni coppia selezionata viene considerata la distanza Jiang-Conrath originale. Una risposta viene classificata come *exact* quando tale distanza è pari a zero entro una tolleranza numerica pari a 0,000001. Una distanza superiore a tale soglia segnala che il concetto scelto non coincide con quello atteso. La tolleranza evita che piccole approssimazioni nei calcoli in virgola mobile influenzino la classificazione.

Quando il numero di risposte non coincide con quello degli elementi corretti, la matrice di matching viene completata con elementi fittizi, ai quali è assegnata una penalità pari a 1,5. Una risposta abbinata a uno di essi viene classificata come *extra*, mentre un elemento corretto privo di una corrispondenza reale viene classificato come *missing*. Il valore 1,5 rappresenta una penalità intermedia rispetto alle distanze Jiang-Conrath valide, normalmente comprese tra 0 e circa 2, e riduce il rischio di forzare associazioni semanticamente poco significative.

Il modulo gestisce anche i casi limite in cui una delle due liste sia vuota. Se lo studente non inserisce risposte per una categoria, tutti gli elementi corretti attesi vengono classificati come *missing*. Se invece sono presenti risposte dello studente ma nel segmento non sono previsti elementi corretti corrispondenti, tali risposte vengono classificate come *extra*.

La stessa logica generale viene applicata in due modi diversi, in base alla struttura del nodo: valutazione a concetto singolo per le *work action* e valutazione strutturata per *data enquiry* e *pharmacological prescription*.

4.5.1 Valutazione dei nodi a concetto singolo

Le *work action* sono valutate come nodi a concetto singolo. Ogni elemento corretto e ogni risposta dello studente corrispondono a un solo concetto finale SNOMED CT. Dopo il matching, la classificazione dipende dalla distanza originale della coppia selezionata.

Se la distanza è pari a zero entro la tolleranza numerica prevista, la risposta viene classificata come *exact*. Se la distanza è superiore alla soglia, viene classificata come *wrong*. Per questo tipo di nodo non viene usata la categoria *partial*, perché non esistono parametri interni che possano essere riconosciuti solo in parte.

Il risultato della valutazione delle *work action* comprende *exact*, *wrong*, *missing* ed *extra*. Queste informazioni vengono successivamente utilizzate per il calcolo dello score e per il feedback mostrato allo studente.

4.5.2 Valutazione dei nodi composti

Data enquiry e *pharmacological prescription* sono valutate come nodi composti. Ogni elemento viene gestito come gruppo di path SNOMED CT, perché il confronto deve mantenere l'appartenenza dei parametri alla relativa richiesta dati o prescrizione farmacologica.

Nel caso delle *data enquiry*, un gruppo può contenere più parametri SNOMED CT. Nel caso delle *pharmacological prescription*, il progetto corrente recupera il riferimento corretto dal nodo come path SNOMED, anche quando la prescrizione fornisce un solo path, essa viene comunque mantenuta come gruppo separato. Questa scelta impedisce di mescolare prescrizioni diverse durante il matching e mantiene uniforme il trattamento degli elementi strutturati.

La valutazione dei nodi composti si articola in due livelli. Nel confronto interno, i parametri di un gruppo dello studente vengono associati a quelli di un gruppo corretto. Da questa fase derivano i match tra parametri, gli elementi missing ed extra e le distanze associate alle corrispondenze reali.

Il risultato del confronto interno viene trasformato in una distanza normalizzata tra gruppi. La somma delle distanze dei parametri associati e delle penalità di 1,5 per gli elementi missing ed extra viene divisa per il numero complessivo degli assegnamenti interni, comprendendo match reali, missing ed extra:

$$distanzaGruppo = \frac{\sum \text{distanza}(\text{parametri associati}) + 1,5 \times \text{missing} + 1,5 \times \text{extra}}{\text{match} + \text{missing} + \text{extra}}$$

I valori normalizzati formano la matrice dei costi utilizzata dall'algoritmo Hungarian per associare tra loro i gruppi completi dello studente e quelli corretti. Le celle dummy eventualmente aggiunte a questa matrice mantengono la penalità pari a 1,5.

La classificazione del gruppo dipende dal risultato interno dei parametri. Un gruppo è exact quando contiene almeno un parametro exact e non presenta parametri wrong, missing o extra. È partial quando contiene almeno un parametro exact, ma presenta anche uno o più problemi interni. È wrong quando nessun parametro interno risulta exact.

La categoria partial è specifica dei nodi composti. Essa rappresenta una risposta nella quale una parte della struttura è stata riconosciuta correttamente, ma l'elemento non coincide interamente con la soluzione attesa. Anche per i nodi composti possono comparire gruppi missing ed extra, secondo lo stesso principio applicato ai nodi a concetto singolo.

4.6 Calcolo dello score didattico e feedback allo studente

Dopo la fase di matching e classificazione, il modulo calcola uno score didattico normalizzato nell'intervallo [0,1], che riassume la qualità complessiva delle risposte fornite dallo studente.

Per i nodi a concetto singolo, come le *work action*, contribuiscono al punteggio solo le corrispondenze classificate come exact. Il denominatore include le associazioni reali prodotte dal matching, gli elementi mancanti e le risposte extra inserite dallo studente. Lo score è definito come:

$$score = \frac{n_{exact}}{n_{match} + n_{missing} + n_{extra}}$$

dove n_{exact} indica il numero di corrispondenze esatte, n_{match} il numero di associazioni reali prodotte dal matching, $n_{missing}$ il numero di elementi corretti non coperti da alcuna risposta e n_{extra} il numero di risposte aggiuntive.

Per i nodi composti, il calcolo dello score avviene su due livelli. Il primo livello riguarda i parametri interni del gruppo. I parametri selezionati dallo studente vengono confrontati con i parametri corretti tramite lo stesso meccanismo usato per i concetti singoli. Lo score interno misura la correttezza dei parametri del gruppo, tenendo conto anche di eventuali parametri mancanti o aggiuntivi.

Il secondo livello riguarda il gruppo nel suo insieme. Dopo il matching tra gruppi, ogni *data enquiry* o *pharmacological prescription* può essere classificata come exact, partial o wrong. Un gruppo exact contribuisce con valore 1, un gruppo wrong contribuisce con valore 0, mentre un gruppo partial contribuisce con lo score interno dei suoi parametri.

Sia m il numero di gruppi associati dal matching, n_{missing} il numero di gruppi corretti non coperti da alcuna risposta e n_{extra} il numero di gruppi aggiuntivi inseriti dallo studente. Lo score dei nodi composti è definito come:

$$\text{score} = \frac{\sum_{i=1}^m c_i}{m + n_{\text{missing}} + n_{\text{extra}}}$$

dove c_i rappresenta il contributo del gruppo associato i . In particolare:

$$c_i = 1, \text{ se il gruppo } i \text{ è exact}$$

$$c_i = s_i, \text{ se il gruppo } i \text{ è partial}$$

$$c_i = 0, \text{ se il gruppo } i \text{ è wrong}$$

dove s_i indica lo score interno del gruppo i , calcolato come il rapporto tra il numero di parametri classificati come exact e il numero complessivo di associazioni interne, parametri missing e parametri extra:

$$s_i = \frac{n_{\text{exact}}^{(i)}}{n_{\text{match}}^{(i)} + n_{\text{missing}}^{(i)} + n_{\text{extra}}^{(i)}}$$

Questa impostazione permette allo score di considerare sia gli elementi correttamente riconosciuti sia le omissioni e gli elementi aggiuntivi. Lo studente riceve punteggio per le parti corrette della risposta, ma viene penalizzato quando omette elementi attesi o inserisce risposte non previste dalla soluzione.

Il feedback mostrato allo studente include il riepilogo delle categorie, lo score e il dettaglio delle associazioni prodotte dal matching. Per le *work action* viene mostrata la coppia risposta-soluzione quando esiste un match reale. Per *data enquiry* e *pharmacological prescription* vengono mostrati anche i parametri interni, così da distinguere l'errore sul gruppo dall'errore su uno o più componenti.

Nei casi wrong o partial, il modulo può mostrare un dettaglio SNOMED. Tale vista mette a confronto il path selezionato dallo studente con il path corretto e permette di osservare la posizione dei concetti nella gerarchia. Il feedback non si limita al punteggio finale, ma rende visibile il motivo della valutazione ricevuta.

4.7 Integrazione nel flusso di GLARE-Edu e ottimizzazione dell'esecuzione

Il modulo di verifica semantica è stato integrato nel flusso di esercitazione già previsto da GLARE-Edu. La valutazione utilizza il pathway associato al case study, i nodi della guideline e i riferimenti SNOMED CT contenuti nei relativi attributi. La suddivisione in segmenti, l'estrazione delle soluzioni corrette e la visualizzazione dei risultati operano direttamente sugli elementi già presenti nell'ambiente.

Per ogni segmento, lo studente inserisce separatamente *work action*, *data enquiry* e *pharmacological prescription*. Dopo l'invio delle risposte, il sistema recupera gli elementi corretti del segmento, esegue il confronto semantico e presenta il risultato prima di passare alla porzione successiva del pathway.

Il calcolo delle distanze richiede l'accesso alle relazioni gerarchiche di SNOMED CT. Al primo caricamento delle risorse semantiche, l'applicazione legge la closure transitiva e costruisce in memoria le strutture necessarie per recuperare gli antenati dei concetti e il numero dei loro discendenti. Nella stessa fase viene acquisito dal database il numero complessivo dei concetti attivi.

Le strutture ottenute vengono mantenute in memoria e riutilizzate durante la valutazione dei segmenti successivi. La closure e il numero totale dei concetti non vengono rilette per ogni coppia confrontata. Questo evita accessi ripetuti alle risorse e riduce il lavoro necessario per il calcolo dell'Information Content, del Lowest Common Subsumer e della distanza Jiang-Conrath.

Durante la valutazione di un segmento, i controlli usati per inserire le risposte vengono temporaneamente disabilitati, così da impedire modifiche mentre il confronto è in corso. Al termine dell'elaborazione, il sistema mostra lo score, le categorie assegnate e il dettaglio delle associazioni individuate dal matching.

L'integrazione riguarda anche la visualizzazione del pathway. All'inizio dell'esercitazione il percorso corretto non viene mostrato, poiché la sua visualizzazione anticiperebbe la soluzione dei segmenti successivi. Dopo la valutazione di un segmento, il sistema rende visibili i nodi e gli archi appartenenti alla porzione corrispondente del percorso corretto.

In presenza di una decisione, il ramo seguito dal case study viene rappresentato come percorso attivo, mentre i rami non percorsi possono essere visualizzati come alternative non attive. La rivelazione progressiva permette allo studente di collegare il feedback ricevuto alla posizione degli elementi nel pathway, mantenendo nascosta la parte del percorso che non è stata ancora verificata.

Il flusso rimane organizzato per segmenti: lo studente inserisce le proprie scelte, riceve la valutazione semantica e osserva la parte del pathway appena verificata. Il riutilizzo delle strutture già caricate in memoria limita le operazioni ripetitive e rende il confronto compatibile con l'interazione richiesta dall'ambiente didattico.

5 Esempi di funzionamento e risultati

In questo capitolo vengono analizzati due segmenti completi della verifica semantica. Per ciascun segmento sono riportate le risposte inserite dallo studente, le soluzioni previste dalla guideline, le distanze calcolate dal modulo, gli assegnamenti prodotti dall'algoritmo Hungarian e lo score finale.

Gli esempi comprendono situazioni differenti: risposte esatte, concetti semanticamente vicini ma non coincidenti, elementi mancanti, risposte aggiuntive e nodi composti con parametri interni. Sono considerate sia le *work action*, valutate come concetti singoli, sia le *data enquiry* e le *pharmacological prescription*, trattate come gruppi strutturati.

I calcoli vengono affiancati alle relative schermate e ai risultati mostrati dall'applicazione. In questo modo è possibile seguire il passaggio dalla risposta inserita dallo studente fino alla classificazione e al feedback mostrati nell'interfaccia.

5.1 Impostazione e criteri di lettura degli esempi

Gli esempi riguardano due segmenti consecutivi dello stesso percorso e sono stati scelti per mostrare il comportamento del modulo in presenza di nodi a concetto singolo e nodi composti.

Per ciascun segmento vengono riportati gli elementi inseriti dallo studente, quelli previsti dalla guideline, le distanze semantiche, gli assegnamenti prodotti dall'algoritmo Hungarian e il risultato finale. Gli esempi raccolgono le principali situazioni utili a verificarne il funzionamento.

Nelle tabelle e nei calcoli, S_i indica una risposta o un gruppo inserito dallo studente, mentre C_j indica un elemento o un gruppo corretto. Per i nodi composti, $S_{\{i,k\}}$ rappresenta il parametro k del gruppo S_i , mentre $C_{\{j,h\}}$ rappresenta il parametro h del gruppo C_j .

Le distanze riportate nelle matrici sono calcolate mediante la misura Jiang-Conrath. L'algoritmo Hungarian utilizza questi valori per individuare l'assegnamento con distanza complessiva minima. Quando il numero degli elementi da confrontare è diverso, la matrice viene completata con elementi dummy aventi penalità pari a 1,5 descritta nel Capitolo 4. In base al lato della matrice in cui compare il dummy, l'elemento non associato viene classificato come missing oppure extra.

Per i nodi composti il matching viene eseguito su due livelli. Per ciascuna possibile coppia di gruppi, l'algoritmo Hungarian associa i parametri interni. La somma delle distanze selezionate e delle penalità per gli elementi missing ed extra viene quindi normalizzata dividendola per il numero totale degli assegnamenti interni. I valori normalizzati formano la matrice utilizzata per stabilire l'assegnamento tra i gruppi dello studente e quelli della guideline.

La distanza semantica viene utilizzata per scegliere le associazioni, mentre la classificazione dipende dal risultato delle singole corrispondenze. Una coppia è *exact* soltanto quando la distanza è pari a zero entro la tolleranza prevista. Nei nodi a concetto singolo, una distanza maggiore di zero produce una classificazione *wrong*. Nei nodi composti, un gruppo viene classificato come *partial* quando contiene almeno un parametro *exact* insieme a uno o più parametri *wrong*, *missing* o *extra*.

5.2 Analisi del Segmento 1

Nel primo segmento la guideline prevede due work action e due data enquiry. Lo studente inserisce entrambe le work action corrette e aggiunge una terza risposta. Per le data enquiry crea due gruppi: il primo contiene un parametro corretto, un parametro diverso da quello atteso e un parametro aggiuntivo, il secondo contiene un solo parametro, che non coincide con nessuno dei tre previsti nel gruppo corretto. Il segmento mostra così la gestione delle corrispondenze *exact*, degli elementi *wrong*, *missing* ed *extra* e la loro influenza sulla classificazione dei gruppi e sullo score.

5.2.1 Work action

La guideline prevede *Haematopoietic stem cell mobilisation* e *Coagulation*. Lo studente seleziona entrambi i concetti e aggiunge *Medical therapy*. La tabella seguente riporta i concetti e i rispettivi path SNOMED CT.

Tipo	ID	Concetto	Path
Corretta	C0	Haematopoietic stem cell mobilisation	Procedure → Stem cell mobilization → Haematopoietic stem cell mobilisation
Corretta	C1	Coagulation	Procedure → Coagulation
Studente	S0	Haematopoietic stem cell mobilisation	Procedure → Stem cell mobilization → Haematopoietic stem cell mobilisation
Studente	S1	Coagulation	Procedure → Coagulation
Studente	S2	Medical therapy	Procedure → Medical therapies

Tabella 1. Concetti e path SNOMED CT delle work action del Segmento 1.

Concetto	descendants	IC
879913000 Haematopoietic stem cell mobilisation	0	$IC = 1 - \log(0 + 1) / \log(372406) = 1.0$
119266004 Coagulation	137	$IC = 1 - \log(137 + 1) / \log(372406) = 0.616$
243121000 Medical therapy	55	$IC = 1 - \log(55 + 1) / \log(372406) = 0.686$
71388002 Procedure	59463	$IC = 1 - \log(59463 + 1) / \log(372406) = 0.143$

Tabella 2. Valori di Information Content delle work action del Segmento 1.

	C0 Stem	C1 Coagulation	Dummy
S0 Stem	$LCS = 879913000$ $JC = 1.0 + 1.0 - 2*1.0 = 0.0$	$LCS = 71388002$ $JC = 1.0 + 0.616 - 2*0.143 = 1.330$	Dummy = 1.5
S1 Coagulation	$LCS = 71388002$ $JC = 0.616 + 1.0 - 2*0.143 = 1.330$	$LCS = 119266004$ $JC = 0.616 + 0.616 - 2*0.616 = 0.0$	Dummy = 1.5
S2 Medical	$LCS = 71388002$ $JC = 0.686 + 1.0 - 2*0.143 = 1.4$	$LCS = 71388002$ $JC = 0.686 + 0.616 - 2*0.143 = 1.016$	Dummy = 1.5

Tabella 3. Matrice delle distanze Jiang-Conrath delle work action del Segmento 1.

Assegnamento	Distanza
S0 Stem → C0 Stem	0.0
S1 Coagulation → C1 Coagulation	0.0
S2 Medical therapy → Dummy	1.5

Tabella 4. Assegnamento delle work action del Segmento 1.

```

=====
SEGMENTO 1 - WORK ACTION
=====

Calcolo delle distanze Jiang-Conrath

S0 Haematopoietic stem cell mobilisation [879913000] -> C0 Haematopoietic stem cell mobilisation [879913000]
concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC studente = 1.000
concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC corretto = 1.000
LCS = 879913000; descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC LCS = 1.000
JC = 1.000 + 1.000 - 2 * 1.000 = 0.000

S0 Haematopoietic stem cell mobilisation [879913000] -> C1 Coagulation [119266004]
concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC studente = 1.000
concetto corretto: descendants = 137; IC = 1 - log(137 + 1) / log(372406)
IC corretto = 0.616
LCS = 71388002; descendants = 59463; IC = 1 - log(59463 + 1) / log(372406)
IC LCS = 0.143
JC = 1.000 + 0.616 - 2 * 0.143 = 1.330

```

Figura 3. Calcolo dell'Information Content e della distanza Jiang-Conrath per S0 rispetto ai nodi corretti nel Segmento 1.

```

S1 Coagulation [119266004] -> C0 Haematopoietic stem cell mobilisation [879913000]
concetto studente: descendants = 137; IC = 1 - log(137 + 1) / log(372406)
IC studente = 0.616
concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC corretto = 1.000
LCS = 71388002; descendants = 59463; IC = 1 - log(59463 + 1) / log(372406)
IC LCS = 0.143
JC = 0.616 + 1.000 - 2 * 0.143 = 1.330

S1 Coagulation [119266004] -> C1 Coagulation [119266004]
concetto studente: descendants = 137; IC = 1 - log(137 + 1) / log(372406)
IC studente = 0.616
concetto corretto: descendants = 137; IC = 1 - log(137 + 1) / log(372406)
IC corretto = 0.616
LCS = 119266004; descendants = 137; IC = 1 - log(137 + 1) / log(372406)
IC LCS = 0.616
JC = 0.616 + 0.616 - 2 * 0.616 = 0.000

S2 Medical therapy [243121000] -> C0 Haematopoietic stem cell mobilisation [879913000]
concetto studente: descendants = 55; IC = 1 - log(55 + 1) / log(372406)
IC studente = 0.686
concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC corretto = 1.000
LCS = 71388002; descendants = 59463; IC = 1 - log(59463 + 1) / log(372406)
IC LCS = 0.143
JC = 0.686 + 1.000 - 2 * 0.143 = 1.400

S2 Medical therapy [243121000] -> C1 Coagulation [119266004]
concetto studente: descendants = 55; IC = 1 - log(55 + 1) / log(372406)
IC studente = 0.686
concetto corretto: descendants = 137; IC = 1 - log(137 + 1) / log(372406)
IC corretto = 0.616
LCS = 71388002; descendants = 59463; IC = 1 - log(59463 + 1) / log(372406)
IC LCS = 0.143
JC = 0.686 + 0.616 - 2 * 0.143 = 1.016

```

Figura 4. Calcolo dell'Information Content e della distanza Jiang-Conrath per S1 e S2 rispetto ai nodi corretti nel Segmento 1.

L'assegnamento globale conserva le due corrispondenze con distanza pari a zero e associa *Medical therapy* al dummy. La risposta aggiuntiva non sostituisce una delle due soluzioni corrette, perché ciò eliminerebbe un match exact e aumenterebbe la somma complessiva delle distanze.

La figura seguente mostra la matrice utilizzata per l'assegnamento e le associazioni selezionate dall'applicazione.

```
Matrice quadrata passata a Hungarian dentro SemanticEvaluator
S0 Haematopoietic stem cell... [879913000] C0 Haematopoietic stem cell... [879913000] C1 Coagulation [119266004] DUMMY-CORRECT
S1 Coagulation [119266004] 0.000 1.330 1.500
S2 Medical therapy [243121000] 1.330 0.000 1.500
S2 Medical therapy [243121000] 1.400 1.016 1.500
Assegnamento prodotto dentro SemanticEvaluator
S0 Haematopoietic stem cell mobilisation [879913000] -> C0 Haematopoietic stem cell mobilisation [879913000] | 0.000
S1 Coagulation [119266004] -> C1 Coagulation [119266004] | 0.000
S2 Medical therapy [243121000] -> DUMMY-CORRECT | 1.500
Totale = 1.500
Risultato finale restituito dall'app
match=2, missing=0, extra=1
score=0.667
```

Figura 5. Matrice delle distanze e assegnamento delle work action del Segmento 1.

Il risultato finale è composto da 2 exact, 0 missing e 1 extra. Poiché lo score include anche gli elementi aggiuntivi nel denominatore, il punteggio delle work action è $2 / (2 + 0 + 1) = 0,667$.

Il caso evidenzia il carattere globale dell'assegnamento. *Medical therapy* presenta una distanza reale rispetto ai concetti corretti, ma entrambe le soluzioni sono già coperte da associazioni con distanza pari a zero. Sostituire uno di questi match con *Medical therapy* aumenterebbe la somma complessiva, l'elemento aggiuntivo viene perciò associato al dummy. La penalizzazione nello score rende inoltre visibile che la risposta non è del tutto precisa, pur contenendo tutte le work action attese.

5.2.2 Data enquiry

Nel Segmento 1 lo studente inserisce due gruppi, S0 e S1, mentre la soluzione contiene i gruppi corretti C0 Query5 e C1 Query1. S0 contiene *Rinse, New patient screening - no abnormality detected* ed *Environment contains animals*, S1 contiene *Asymptomatic*. C0 Query5 contiene *Rinse* e *No consent for MR - Measles/rubella vaccine*, mentre C1 Query1 contiene *Clinical privilege*, *Retigabine* e *Identification of fungus by MALDI-TOF MS*.

Il modulo confronta ogni gruppo dello studente con ogni gruppo corretto: S0-C0, S0-C1, S1-C0 e S1-C1. Per ciascuna coppia, l'algoritmo Hungarian associa i parametri interni, la somma delle distanze selezionate e delle penalità viene poi divisa per il numero totale degli assegnamenti interni. Le distanze normalizzate così ottenute formano la matrice dei costi utilizzata per determinare l'assegnamento complessivo tra i gruppi dello studente e quelli previsti dalla guideline.

Origine e gruppo	ID	Concetto	Path
Studente S0	S0.0	Rinse	Qualifier value → Dose form administration method → Rinse
Studente S0	S0.1	New patient screening - no abnormality detected	Clinical finding → Temporal finding → New patient screening - no abnormality detected
Studente S0	S0.2	Environment contains animals	Clinical finding → Safety finding → Environment contains animals
Studente S1	S1.0	Asymptomatic	Clinical finding → Absence of symptoms
Guideline C0 Query5	C0.0	Rinse	Qualifier value → Dose form administration method → Rinse
Guideline C0 Query5	C0.1	No consent for MR - Measles/rubella vaccine	Clinical finding → Temporal finding → No consent for MR - Measles/rubella vaccine
Guideline C1 Query1	C1.0	Clinical privilege	Clinical finding → Administrative statuses → Clinical privilege
Guideline C1 Query1	C1.1	Retigabine	Substance → Drug or medicament → Central nervous system agent → Anticonvulsant → Retigabine
Guideline C1 Query1	C1.2	Identification of fungus by MALDI-TOF MS	Observable entity → Identification of fungus by MALDI-TOF MS

Tabella 5. Gruppi e parametri delle data enquiry del Segmento 1.

Confronti tra i gruppi

Confronto tra S0 e C0 Query5

Concetto	descendants	IC
782155003 Rinse	0	$IC = 1 - \log(0 + 1) / \log(372406) = 1.0$
171293001 No consent for MR - Measles/rubella vaccine	0	$IC = 1 - \log(0 + 1) / \log(372406) = 1.0$
171327009 New patient screening - no abnormality detected	0	$IC = 1 - \log(0 + 1) / \log(372406) = 1.0$
425338002 Environment contains animals	5	$IC = 1 - \log(5 + 1) / \log(372406) = 0.860$
118226009 Temporal finding	4	$IC = 1 - \log(4 + 1) / \log(372406) = 0.875$
404684003 Clinical finding	124873	$IC = 1 - \log(124873 + 1) / \log(372406) = 0.085$
138875005 SNOMED CT Concept	372405	$IC = 1 - \log(372405 + 1) / \log(372406) = 0.0$

Tabella 6. Valori di Information Content usati nel confronto S0-C0 Query5.

	C0.0 Rinse	C0.1 MR no consent	Dummy
S0.0 Rinse	LCS = 782155003 $JC = 1.0 + 1.0 - 2 \cdot 1.0 = 0.0$	LCS = 138875005 $JC = 1.0 + 1.0 - 2 \cdot 0.0 = 2.0$	Dummy = 1.5
S0.1 New patient	LCS = 138875005 $JC = 1.0 + 1.0 - 2 \cdot 0.0 = 2.0$	LCS = 118226009 $JC = 1.0 + 1.0 - 2 \cdot 0.875 = 0.251$	Dummy = 1.5
S0.2 Animals	LCS = 138875005 $JC = 0.860 + 1.0 - 2 \cdot 0.0 = 1.860$	LCS = 404684003 $JC = 0.860 + 1.0 - 2 \cdot 0.085 = 1.690$	Dummy = 1.5

Tabella 7. Matrice delle distanze interne tra S0 e C0 Query5.

Assegnamento	Distanza
S0.0 Rinse → C0.0 Rinse	0.0
S0.1 New patient screening → C0.1 No consent for MR	0.251
S0.2 Environment contains animals → Dummy	1.5

Tabella 8. Assegnamento interno tra S0 e C0 Query5.

```

=====
SEGMENTO 1 - DATA ENQUIRY
=====

Primo livello: S0 contro C0 Query5

Calcolo delle distanze Jiang-Conrath

S0.0 Rinse [782155003] -> C0.0 Rinse [782155003]
  concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
  IC studente = 1.000
  concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
  IC corretto = 1.000
  LCS = 782155003; descendants = 0; IC = 1 - log(0 + 1) / log(372406)
  IC LCS = 1.000
  JC = 1.000 + 1.000 - 2 * 1.000 = 0.000

S0.0 Rinse [782155003] -> C0.1 No consent for MR [171293001]
  concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
  IC studente = 1.000
  concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
  IC corretto = 1.000
  LCS = 138875005; descendants = 372405; IC = 1 - log(372405 + 1) / log(372406)
  IC LCS = 0.000
  JC = 1.000 + 1.000 - 2 * 0.000 = 2.000

S0.1 New patient screening [171327009] -> C0.0 Rinse [782155003]
  concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
  IC studente = 1.000
  concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
  IC corretto = 1.000
  LCS = 138875005; descendants = 372405; IC = 1 - log(372405 + 1) / log(372406)
  IC LCS = 0.000
  JC = 1.000 + 1.000 - 2 * 0.000 = 2.000

```

Figura 6. Calcolo dell'Information Content e della distanza Jiang-Conrath tra S0 e C0 Query5, prima parte.

```

S0.1 New patient screening [171327009] -> C0.1 No consent for MR [171293001]
  concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
  IC studente = 1.000
  concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
  IC corretto = 1.000
  LCS = 118226009; descendants = 4; IC = 1 - log(4 + 1) / log(372406)
  IC LCS = 0.875
  JC = 1.000 + 1.000 - 2 * 0.875 = 0.251

S0.2 Environment contains animals [425338002] -> C0.0 Rinse [782155003]
  concetto studente: descendants = 5; IC = 1 - log(5 + 1) / log(372406)
  IC studente = 0.860
  concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
  IC corretto = 1.000
  LCS = 138875005; descendants = 372405; IC = 1 - log(372405 + 1) / log(372406)
  IC LCS = 0.000
  JC = 0.860 + 1.000 - 2 * 0.000 = 1.860

S0.2 Environment contains animals [425338002] -> C0.1 No consent for MR [171293001]
  concetto studente: descendants = 5; IC = 1 - log(5 + 1) / log(372406)
  IC studente = 0.860
  concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
  IC corretto = 1.000
  LCS = 404684003; descendants = 124873; IC = 1 - log(124873 + 1) / log(372406)
  IC LCS = 0.085
  JC = 0.860 + 1.000 - 2 * 0.085 = 1.690

```

Figura 7. Calcolo dell'Information Content e della distanza Jiang-Conrath tra S0 e C0 Query5, seconda parte.

L'assegnamento interno associa Rinse al parametro corretto omonimo con distanza pari a zero, New patient screening a No consent for MR con distanza pari a 0,251 ed Environment contains animals al dummy come parametro extra. La somma interna è 1,751, divisa per i tre assegnamenti considerati, produce una distanza normalizzata tra S0 e C0 pari a 0,584.

```
Matrice quadrata passata a Hungarian dentro SemanticEvaluator
S0.0 Rinse [782155003]          C0.0 Rinse [782155003]          C0.1 No consent for MR [171293001]    DUMMY-CORRECT
0.000                          2.000                          1.500
S0.1 New patient screening [171327009]  2.000                          0.251                          1.500
S0.2 Environment contains a... [425338002] 1.860                          1.690                          1.500
Assegnamento prodotto dentro SemanticEvaluator
S0.0 Rinse [782155003] -> C0.0 Rinse [782155003] | 0.000
S0.1 New patient screening [171327009] -> C0.1 No consent for MR [171293001] | 0.251
S0.2 Environment contains animals [425338002] -> DUMMY-CORRECT | 1.500
Totale = 1.751
```

Figura 8. Matrice delle distanze e assegnamento interno tra S0 e C0 Query5.

Confronto tra S0 e C1 Query1

	C1.0 Clinical privilege	C1.1 Retigabine	C1.2 Fungus ID
S0.0 Rinse	2.0	2.0	2.0
S0.1 New patient screening	1.830	2.0	2.0
S0.2 Environment contains animals	1.690	1.860	1.860
Assegnamento	S0.1 → C1.0 = 1.830	S0.0 → C1.1 = 2.0	S0.2 → C1.2 = 1.860

Tabella 9. Matrice e assegnamento interni tra S0 e C1 Query1.

```
Primo livello: S0 contro C1 Query1
Calcolo delle distanze Jiang-Conrath
S0.0 Rinse [782155003] -> C1.0 Clinical privilege [736014004]
concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC studente = 1.000
concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC corretto = 1.000
LCS = 138875005; descendants = 372405; IC = 1 - log(372405 + 1) / log(372406)
IC LCS = 0.000
JC = 1.000 + 1.000 - 2 * 0.000 = 2.000
S0.0 Rinse [782155003] -> C1.1 Retigabine [699271005]
concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC studente = 1.000
concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC corretto = 1.000
LCS = 138875005; descendants = 372405; IC = 1 - log(372405 + 1) / log(372406)
IC LCS = 0.000
JC = 1.000 + 1.000 - 2 * 0.000 = 2.000
S0.0 Rinse [782155003] -> C1.2 Identification of fungus by MALDI-TOF MS [1285667006]
concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC studente = 1.000
concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC corretto = 1.000
LCS = 138875005; descendants = 372405; IC = 1 - log(372405 + 1) / log(372406)
IC LCS = 0.000
JC = 1.000 + 1.000 - 2 * 0.000 = 2.000
```

Figura 9. Calcolo della distanza Jiang-Conrath tra S0.0 e i parametri di C1 Query1.

```

S0.1 New patient screening [171327009] -> C1.0 Clinical privilege [736014004]
concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC studente = 1.000
concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC corretto = 1.000
LCS = 404684003; descendants = 124873; IC = 1 - log(124873 + 1) / log(372406)
IC LCS = 0.085
JC = 1.000 + 1.000 - 2 * 0.085 = 1.830

S0.1 New patient screening [171327009] -> C1.1 Retigabine [699271005]
concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC studente = 1.000
concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC corretto = 1.000
LCS = 138875005; descendants = 372405; IC = 1 - log(372405 + 1) / log(372406)
IC LCS = 0.000
JC = 1.000 + 1.000 - 2 * 0.000 = 2.000

S0.1 New patient screening [171327009] -> C1.2 Identification of fungus by MALDI-TOF MS [1285667006]
concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC studente = 1.000
concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC corretto = 1.000
LCS = 138875005; descendants = 372405; IC = 1 - log(372405 + 1) / log(372406)
IC LCS = 0.000
JC = 1.000 + 1.000 - 2 * 0.000 = 2.000

```

Figura 10. Calcolo della distanza Jiang-Conrath tra S0.1 e i parametri di C1 Query1.

```

S0.2 Environment contains animals [425338002] -> C1.0 Clinical privilege [736014004]
concetto studente: descendants = 5; IC = 1 - log(5 + 1) / log(372406)
IC studente = 0.860
concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC corretto = 1.000
LCS = 404684003; descendants = 124873; IC = 1 - log(124873 + 1) / log(372406)
IC LCS = 0.085
JC = 0.860 + 1.000 - 2 * 0.085 = 1.690

S0.2 Environment contains animals [425338002] -> C1.1 Retigabine [699271005]
concetto studente: descendants = 5; IC = 1 - log(5 + 1) / log(372406)
IC studente = 0.860
concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC corretto = 1.000
LCS = 138875005; descendants = 372405; IC = 1 - log(372405 + 1) / log(372406)
IC LCS = 0.000
JC = 0.860 + 1.000 - 2 * 0.000 = 1.860

S0.2 Environment contains animals [425338002] -> C1.2 Identification of fungus by MALDI-TOF MS [1285667006]
concetto studente: descendants = 5; IC = 1 - log(5 + 1) / log(372406)
IC studente = 0.860
concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC corretto = 1.000
LCS = 138875005; descendants = 372405; IC = 1 - log(372405 + 1) / log(372406)
IC LCS = 0.000
JC = 0.860 + 1.000 - 2 * 0.000 = 1.860

```

Figura 11. Calcolo della distanza Jiang-Conrath tra S0.2 e i parametri di C1 Query1.

```

Matrice quadrata passata a Hungarian dentro SemanticEvaluator
          C1.0 Clinical privilege [736014004]      C1.1 Retigabine [699271005]      C1.2 Identification of fungus by MALDI-TOF MS [1285667006]
S0.0 Rinse [782155003]      2.000      2.000      2.000
S0.1 New patient screening [171327009]      1.830      2.000      2.000
S0.2 Environment contains animals [425338002]      1.690      1.860      1.860
Assegnamento prodotto dentro SemanticEvaluator
S0.0 Rinse [782155003] -> C1.1 Retigabine [699271005] | 2.000
S0.1 New patient screening [171327009] -> C1.0 Clinical privilege [736014004] | 1.830
S0.2 Environment contains animals [425338002] -> C1.2 Identification of fungus by MALDI-TOF MS [1285667006] | 1.860
Totale = 5.690

```

Figura 12. Matrice delle distanze e assegnamento interno tra S0 e C1 Query1.

L'assegnamento interno associa S0.0 a C1.1 con distanza 2,000, S0.1 a C1.0 con distanza 1,830 e S0.2 a C1.2 con distanza 1,860. La somma interna è 5,690, divisa per i tre assegnamenti considerati, produce una distanza normalizzata tra S0 e C1 pari a 1,897.

Confronto tra S1 e C0 Query5

	C0.0 Rinse	C0.1 No consent for MR
S1.0 Asymptomatic	1.673	1.503
DUMMY-STUDENT	1.5	1.5
Assegnamento	DUMMY-STUDENT → C0.0 = 1.5	S1.0 → C0.1 = 1.503

Tabella 10. Matrice e assegnamento interni tra S1 e C0 Query5.

```
Primo livello: S1 contro C0 Query5

Calcolo delle distanze Jiang-Conrath

S1.0 Asymptomatic [84387000] -> C0.0 Rinse [782155003]
  concetto studente: descendants = 65; IC = 1 - log(65 + 1) / log(372406)
  IC studente = 0.673
  concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
  IC corretto = 1.000
  LCS = 138875005; descendants = 372405; IC = 1 - log(372405 + 1) / log(372406)
  IC LCS = 0.000
  JC = 0.673 + 1.000 - 2 * 0.000 = 1.673

S1.0 Asymptomatic [84387000] -> C0.1 No consent for MR [171293001]
  concetto studente: descendants = 65; IC = 1 - log(65 + 1) / log(372406)
  IC studente = 0.673
  concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
  IC corretto = 1.000
  LCS = 404684003; descendants = 124873; IC = 1 - log(124873 + 1) / log(372406)
  IC LCS = 0.085
  JC = 0.673 + 1.000 - 2 * 0.085 = 1.503
```

Figura 13. Calcolo della distanza Jiang-Conrath tra S1 e C0 Query5.

```
Matrice quadrata passata a Hungarian dentro SemanticEvaluator
S1.0 Asymptomatic [84387000]      C0.0 Rinse [782155003]      C0.1 No consent for MR
DUMMY-STUDENT                     1.673                     1.503
Assegnamento prodotto dentro SemanticEvaluator
S1.0 Asymptomatic [84387000] -> C0.1 No consent for MR [171293001] | 1.503
DUMMY-STUDENT -> C0.0 Rinse [782155003] | 1.500
Totale = 3.003
```

Figura 14. Matrice delle distanze e assegnamento interno tra S1 e C0 Query5.

Il concetto Asymptomatic viene associato a No consent for MR con distanza pari a 1,503, mentre Rinse rimane missing e contribuisce con la penalità dummy di 1,5. La somma interna è 3,003, divisa per i due assegnamenti considerati, produce una distanza normalizzata tra S1 e C0 pari a 1,502.

Confronto tra S1 e C1 Query1

Concetto	descendants	IC
84387000 Asymptomatic	65	$IC = 1 - \log(65 + 1) / \log(372406) = 0.673$
736014004 Clinical privilege	0	$IC = 1 - \log(0 + 1) / \log(372406) = 1.0$
699271005 Retigabine	0	$IC = 1 - \log(0 + 1) / \log(372406) = 1.0$
1285667006 Identification of fungus by MALDI-TOF MS	0	$IC = 1 - \log(0 + 1) / \log(372406) = 1.0$
404684003 Clinical finding	124873	$IC = 1 - \log(124873 + 1) / \log(372406) = 0.085$
138875005 SNOMED CT Concept	372405	$IC = 1 - \log(372405 + 1) / \log(372406) = 0.0$

Tabella 11. Valori di Information Content usati nel confronto S1-C1 Query1.

	C1.0 Clinical privilege	C1.1 Retigabine	C1.2 Fungus ID
S1.0 Asymptomatic	LCS =404684003 $JC = 0.673 + 1.0 - 2*0.085 = 1.503$	LCS =138875005 $JC = 0.673 + 1.0 - 2*0.0 = 1.673$	LCS =138875005 $JC = 0.673 + 1.0 - 2*0.0 = 1.673$
Dummy	Dummy = 1.5	Dummy = 1.5	Dummy = 1.5
Dummy	Dummy = 1.5	Dummy = 1.5	Dummy = 1.5

Tabella 12. Matrice delle distanze interne tra S1 e C1 Query1.

Assegnamento	Distanza
S1.0 Asymptomatic → C1.0 Clinical privilege	1.503
Dummy → C1.1 Retigabine	1.5
Dummy → C1.2 Identification of fungus by MALDI-TOF MS	1.5

Tabella 13. Assegnamento interno tra S1 e C1 Query1.

```

Primo livello: S1 contro C1 Query1

Calcolo delle distanze Jiang-Conrath

S1.0 Asymptomatic [84387000] -> C1.0 Clinical privilege [736014004]
  concetto studente: descendants = 65; IC = 1 - log(65 + 1) / log(372406)
  IC studente = 0.673
  concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
  IC corretto = 1.000
  LCS = 404684003; descendants = 124873; IC = 1 - log(124873 + 1) / log(372406)
  IC LCS = 0.085
  JC = 0.673 + 1.000 - 2 * 0.085 = 1.503

S1.0 Asymptomatic [84387000] -> C1.1 Retigabine [699271005]
  concetto studente: descendants = 65; IC = 1 - log(65 + 1) / log(372406)
  IC studente = 0.673
  concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
  IC corretto = 1.000
  LCS = 138875005; descendants = 372405; IC = 1 - log(372405 + 1) / log(372406)
  IC LCS = 0.000
  JC = 0.673 + 1.000 - 2 * 0.000 = 1.673

S1.0 Asymptomatic [84387000] -> C1.2 Identification of fungus by MALDI-TOF MS [1285667006]
  concetto studente: descendants = 65; IC = 1 - log(65 + 1) / log(372406)
  IC studente = 0.673
  concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
  IC corretto = 1.000
  LCS = 138875005; descendants = 372405; IC = 1 - log(372405 + 1) / log(372406)
  IC LCS = 0.000
  JC = 0.673 + 1.000 - 2 * 0.000 = 1.673

```

Figura 15. Calcolo dell'Information Content e della distanza Jiang-Conrath tra S1 e C1 Query1.

```

Matrice quadrata passata a Hungarian dentro SemanticEvaluator
      C1.0 Clinical privilege [736014004]    C1.1 Retigabine [699271005]    C1.2 Identification of fungus by MALDI-TOF MS [1285667006]
S1.0 Asymptomatic [84387000]              1.503                        1.673                        1.673
DUMMY-STUDENT                             1.500                        1.500                        1.500
DUMMY-STUDENT                             1.500                        1.500                        1.500
Assegnamento prodotto dentro SemanticEvaluator
S1.0 Asymptomatic [84387000] -> C1.0 Clinical privilege [736014004] | 1.503
DUMMY-STUDENT -> C1.1 Retigabine [699271005] | 1.500
DUMMY-STUDENT -> C1.2 Identification of fungus by MALDI-TOF MS [1285667006] | 1.500
Totale = 4.503

```

Figura 16. Matrice delle distanze e assegnamento interno tra S1 e C1 Query1.

Il concetto Asymptomatic viene associato a Clinical privilege con distanza pari a 1,503. Retigabine e Identification of fungus by MALDI-TOF MS rimangono missing. La somma interna è 4,503. divisa per i tre assegnamenti considerati, produce una distanza normalizzata tra S1 e C1 pari a 1,501.

Assegnamento finale tra le data enquiry

Ogni cella della matrice di assegnamento tra gruppi rappresenta la distanza normalizzata tra un gruppo dello studente e un gruppo corretto. Al numeratore vengono sommate le distanze dei parametri associati e le penalità per gli elementi interni missing o extra, il risultato viene diviso per il numero totale degli assegnamenti interni.

$$distanza_{gruppo}(Sx, Cy) = \frac{\text{somma delle distanze interne} + 1,5 \cdot \text{missing} + 1,5 \cdot \text{extra}}{\text{match} + \text{missing} + \text{extra}}$$

Le quattro celle della matrice sono ottenute dai confronti interni riportati nelle tabelle precedenti:

$$distanza_gruppo(\text{nodo S0}, \text{nodo C0 Query5}) = (0,000 + 0,251 + 1,500) / 3 = 0,584$$

$$distanza_gruppo(\text{nodo S0}, \text{nodo C1 Query1}) = (2,000 + 1,830 + 1,860) / 3 = 1,897$$

$$distanza_gruppo(\text{nodo S1}, \text{nodo C0 Query5}) = (1,503 + 1,500) / 2 = 1,502$$

$$distanza_gruppo(\text{nodo S1}, \text{nodo C1 Query1}) = (1,503 + 1,500 + 1,500) / 3 = 1,501$$

Cella	Somma delle distanze e penalità	Assegnamenti considerati	Valore normalizzato
nodo S0 → nodo C0 Query5	1.751	3	0.584
nodo S0 → nodo C1 Query1	5.690	3	1.897
nodo S1 → nodo C0 Query5	3.003	2	1.502
nodo S1 → nodo C1 Query1	4.503	3	1.501

Tabella 14. Distanze normalizzate tra i gruppi di data enquiry del Segmento 1.

Gruppo studente	nodo C0 Query5	nodo C1 Query1
nodo S0	0.584	1.897
nodo S1	1.502	1.501

Tabella 15. Matrice di assegnamento tra i gruppi delle data enquiry del Segmento 1.

Possibilità	Calcolo	Totale
A	nodo S0 → nodo C0 Query5 + nodo S1 → nodo C1 Query1 = 0.584 + 1.501	2.085
B	nodo S0 → nodo C1 Query1 + nodo S1 → nodo C0 Query5 = 1.897 + 1.502	3.399

Tabella 16. Confronto tra i possibili assegnamenti delle data enquiry del Segmento 1.

L'assegnamento S0 → C0 Query5 e S1 → C1 Query1 produce un costo pari a 2,085, inferiore al valore 3,399 dell'associazione incrociata. L'algoritmo seleziona pertanto il primo abbinamento.

La figura seguente conferma sia i valori delle quattro celle sia l'assegnamento finale tra i gruppi.

```

          C0 Query5      C1 Query1
S0          0.584        1.897
S1          1.502        1.501

CONFRONTO ASSEGNAMENTI

A  S0 -> C0 Query5  +  S1 -> C1 Query1
    0.584 + 1.501 = 2.085    SCELTA

B  S0 -> C1 Query1  +  S1 -> C0 Query5
    1.897 + 1.502 = 3.399

RISULTATO MOSTRATO ALLO STUDENTE

exact = 0    partial = 1    wrong = 1    missing = 0    extra = 0
score = 0.167

```

Figura 17. Matrice e assegnamento tra i gruppi delle data enquiry del Segmento 1.

Il gruppo S0 è partial e contribuisce con il proprio score interno, il gruppo S1 è wrong e contribuisce con valore zero. Lo score complessivo delle data enquiry del Segmento 1 è 0,167.

Nel complesso, il Segmento 1 separa chiaramente tre aspetti della valutazione. Le distanze stabiliscono gli abbinamenti, le categorie descrivono la qualità delle singole corrispondenze e lo score sintetizza il risultato tenendo conto anche degli elementi extra.

5.3 Analisi del Segmento 2

Nel secondo segmento la guideline prevede tre work action, due data enquiry composte rispettivamente da tre e due parametri e due pharmacological prescription separate, ciascuna contenente un solo concetto. Lo studente inserisce due work action semanticamente vicine ma non coincidenti, una sola data enquiry con un unico parametro e due gruppi di pharmacological prescription: nel primo riunisce i due concetti attesi, mentre nel secondo inserisce un concetto semanticamente vicino a uno di essi.

5.3.1 Work action

La guideline contiene *Emergency procedure*, *Dysphagia therapy regime* e *Application of hip spica cast*. Lo studente seleziona *Chest thrust* e *Galvanism*. Le due risposte vengono associate alle soluzioni che minimizzano la distanza, mentre la terza work action corretta rimane missing.

Tipo	ID	Concetto	Path
Corretta	C0	Emergency procedure	Procedure → Emergency procedure
Corretta	C1	Dysphagia therapy regime	Procedure → Regimes and therapies → Dysphagia therapy regime
Corretta	C2	Application of hip spica cast	Procedure → Application of hip spica cast

Tipo	ID	Concetto	Path
Studente	S0	Chest thrust	Procedure → Emergency procedure → First aid → Chest thrust
Studente	S1	Galvanism	Procedure → Galvanism

Tabella 17. Concetti e path SNOMED CT delle work action del Segmento 2.

Concetto	descendants	IC
703985001 Chest thrust	0	$IC = 1 - \log(0 + 1) / \log(372406) = 1.0$
773262000 Galvanism	0	$IC = 1 - \log(0+1) / \log(372406) = 1.0$
373110003 Emergency procedure	141	$IC = 1 - \log(141 + 1) / \log(372406) = 0.614$
311569007 Dysphagia therapy regime	9	$IC = 1 - \log(9 + 1) / \log(372406) = 0.820$
127599007 Application of hip spica cast	3	$IC = 1 - \log(3 + 1) / \log(372406) = 0.892$
71388002 Procedure	59463	$IC = 1 - \log(59463 + 1) / \log(372406) = 0.143$

Tabella 18. Valori di Information Content delle work action del Segmento 2.

	C0 Emergency	C1 Dysphagia	C2 Hip spica cast
S0 Chest	LCS = 373110003 JC = $1.0 + 0.614 - 2*0.614 = 0.386$	LCS = 71388002 JC = $1.0 + 0.820 - 2*0.143 = 1.534$	LCS = 71388002 JC = $1.0 + 0.892 - 2*0.143 = 1.606$
S1 Galvanism	LCS = 71388002 JC = $1.0 + 0.614 - 2*0.143 = 1.328$	LCS = 71388002 JC = $1.0 + 0.820 - 2*0.143 = 1.534$	LCS = 71388002 JC = $1.0 + 0.892 - 2*0.143 = 1.606$
Dummy	Dummy = 1.5	Dummy = 1.5	Dummy = 1.5

Tabella 19. Matrice delle distanze Jiang-Conrath delle work action del Segmento 2.

Assegnamento	Distanza
S0 Chest thrust → C0 Emergency procedure	0.386
S1 Galvanism → C1 Dysphagia therapy regime	1.534
Dummy → C2 Application of hip spica cast	1.5

Tabella 20. Assegnamento delle work action del Segmento 2.

```

=====
SEGMENTO 2 - WORK ACTION
=====

Calcolo delle distanze Jiang-Conrath

S0 Chest thrust [703985001] -> C0 Emergency procedure [373110003]
concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC studente = 1.000
concetto corretto: descendants = 141; IC = 1 - log(141 + 1) / log(372406)
IC corretto = 0.614
LCS = 373110003; descendants = 141; IC = 1 - log(141 + 1) / log(372406)
IC LCS = 0.614
JC = 1.000 + 0.614 - 2 * 0.614 = 0.386

S0 Chest thrust [703985001] -> C1 Dysphagia therapy regime [311569007]
concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC studente = 1.000
concetto corretto: descendants = 9; IC = 1 - log(9 + 1) / log(372406)
IC corretto = 0.820
LCS = 71388002; descendants = 59463; IC = 1 - log(59463 + 1) / log(372406)
IC LCS = 0.143
JC = 1.000 + 0.820 - 2 * 0.143 = 1.534

S0 Chest thrust [703985001] -> C2 Application of hip spica cast [127599007]
concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC studente = 1.000
concetto corretto: descendants = 3; IC = 1 - log(3 + 1) / log(372406)
IC corretto = 0.892
LCS = 71388002; descendants = 59463; IC = 1 - log(59463 + 1) / log(372406)
IC LCS = 0.143
JC = 1.000 + 0.892 - 2 * 0.143 = 1.606

```

Figura 18. Calcolo dell'Information Content e della distanza Jiang-Conrath per S0 rispetto a C0, C1 e C2 nel Segmento 2.

```

S1 Galvanism [773262000] -> C0 Emergency procedure [373110003]
concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC studente = 1.000
concetto corretto: descendants = 141; IC = 1 - log(141 + 1) / log(372406)
IC corretto = 0.614
LCS = 71388002; descendants = 59463; IC = 1 - log(59463 + 1) / log(372406)
IC LCS = 0.143
JC = 1.000 + 0.614 - 2 * 0.143 = 1.328

S1 Galvanism [773262000] -> C1 Dysphagia therapy regime [311569007]
concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC studente = 1.000
concetto corretto: descendants = 9; IC = 1 - log(9 + 1) / log(372406)
IC corretto = 0.820
LCS = 71388002; descendants = 59463; IC = 1 - log(59463 + 1) / log(372406)
IC LCS = 0.143
JC = 1.000 + 0.820 - 2 * 0.143 = 1.534

S1 Galvanism [773262000] -> C2 Application of hip spica cast [127599007]
concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC studente = 1.000
concetto corretto: descendants = 3; IC = 1 - log(3 + 1) / log(372406)
IC corretto = 0.892
LCS = 71388002; descendants = 59463; IC = 1 - log(59463 + 1) / log(372406)
IC LCS = 0.143
JC = 1.000 + 0.892 - 2 * 0.143 = 1.606

```

Figura 19. Calcolo dell'Information Content e della distanza Jiang-Conrath per S1 rispetto a C0, C1 e C2 nel Segmento 2.

L'assegnamento associa *Chest thrust* a *Emergency procedure* con distanza 0,386 e *Galvanism* a *Dysphagia therapy regime* con distanza 1,534. *Application of hip spica cast* viene associata a una riga dummy. Le due coppie reali sono wrong, perché la loro distanza è maggiore di zero, il terzo elemento è missing.

La figura seguente mostra la matrice Jiang-Conrath e l'assegnamento scelto dall'algoritmo.

```
Matrice quadrata passata a Hungarian dentro SemanticEvaluator
      C0 Emergency procedure [373110003]      C1 Dysphagia therapy regime [311569007]      C2 Application of
S0 Chest thrust [703985001]      0.386      1.534      1.606
S1 Galvanism [773262000]      1.328      1.534      1.606
DUMMY-STUDENT      1.500      1.500      1.500
Assegnamento prodotto dentro SemanticEvaluator
S0 Chest thrust [703985001] -> C0 Emergency procedure [373110003] | 0.386
S1 Galvanism [773262000] -> C1 Dysphagia therapy regime [311569007] | 1.534
DUMMY-STUDENT -> C2 Application of hip spica cast [127599007] | 1.500
Totale = 3.421
Risultato finale restituito dall'app
match=2, missing=1, extra=0
score=0.000
```

Figura 20. Matrice delle distanze e assegnamento delle work action del Segmento 2.

5.3.2 Data enquiry

La soluzione contiene due gruppi corretti, C0 Query2 e C1 Query3, mentre lo studente inserisce il solo gruppo S0, composto dal parametro *Concentration of neurotransmitter in blood*. S0 viene confrontato separatamente con ciascuno dei due gruppi corretti applicando l'algoritmo Hungarian ai rispettivi parametri interni. Le distanze complessive ottenute vengono normalizzate rispetto al numero di assegnamenti interni e poi utilizzate per costruire la matrice di assegnamento tra i gruppi.

Poiché il numero dei gruppi inseriti dallo studente è inferiore a quello previsto dalla soluzione, la matrice viene completata con una riga dummy, che rappresenta il gruppo corretto rimasto senza corrispondenza.

Origine e gruppo	ID	Concetto	Path
Corretta C0 Query2	C0.0	Concentration of neurotransmitter in blood	Observable entity → Concentration of neurotransmitter in blood
Corretta C0 Query2	C0.1	Administrative statuses	Clinical finding → Administrative statuses
Corretta C0 Query2	C0.2	Body product observable	Observable entity → Body product observable
Corretta C1 Query3	C1.0	Carrier oil-containing product	Qualifier value → Role → Carrier oil → Carrier oil-containing product
Corretta C1 Query3	C1.1	Oromucosal gel	Qualifier value → Pharmaceutical dose form → Oromucosal dose form → Oromucosal gel

Origine e gruppo	ID	Concetto	Path
Studente S0	S0.0	Concentration of neurotransmitter in blood	Observable entity → Concentration of neurotransmitter in blood

Tabella 21. Concetti e path SNOMED CT dei gruppi di data enquiry del Segmento 2.

Concetto	descendants	IC
1285434004 Concentration of neurotransmitter in blood	0	$IC = 1 - \log(0 + 1) / \log(372406) = 1.0$
307824009 Administrative statuses	1918	$IC = 1 - \log(1918 + 1) / \log(372406) = 0.411$
364684009 Body product observable	62	$IC = 1 - \log(62 + 1) / \log(372406) = 0.677$
838458004 Carrier oil-containing product	0	$IC = 1 - \log(0 + 1) / \log(372406) = 1.0$
385078007 Oromucosal gel	1	$IC = 1 - \log(1 + 1) / \log(372406) = 0.946$
363787002 Observable entity	10861	$IC = 1 - \log(10861 + 1) / \log(372406) = 0.276$
138875005 SNOMED CT Concept	372405	$IC = 1 - \log(372405 + 1) / \log(372406) = 0.0$

Tabella 22. Valori di Information Content usati nei confronti delle data enquiry del Segmento 2.

Confronto tra S0 e C0 Query2

	C0.0 Neuro	C0.1 Administrative	C0.2 Body product
S0.0 Neuro	LCS=1285434004 $JC = 1.0 + 1.0 - 2*1.0 = 0.0$	LCS=138875005 $JC = 1.0 + 0.411 - 2*0.0 = 1.411$	LCS=363787002 $JC = 1.0 + 0.677 - 2*0.276 = 1.126$
Dummy	Dummy = 1.5	Dummy = 1.5	Dummy = 1.5
Dummy	Dummy = 1.5	Dummy = 1.5	Dummy = 1.5

Tabella 23. Matrice delle distanze interne tra S0 e C0 Query2.

Assegnamento	Distanza
S0.0 Neurotransmitter → C0.0 Neurotransmitter	0.0
Dummy → C0.1 Administrative statuses	1.5
Dummy → C0.2 Body product observable	1.5

Tabella 24. Assegnamento interno tra S0 e C0 Query2.

```

=====
SEGMENTO 2 - DATA ENQUIRY
=====

Primo livello: S0 contro C0 Query2

Calcolo delle distanze Jiang-Conrath

S0.0 Concentration of neurotransmitter in blood [1285434004] -> C0.0 Concentration of neurotransmitter in blood [1285434004]
concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC studente = 1.000
concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC corretto = 1.000
LCS = 1285434004; descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC LCS = 1.000
JC = 1.000 + 1.000 - 2 * 1.000 = 0.000

S0.0 Concentration of neurotransmitter in blood [1285434004] -> C0.1 Administrative statuses [307824009]
concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC studente = 1.000
concetto corretto: descendants = 1918; IC = 1 - log(1918 + 1) / log(372406)
IC corretto = 0.411
LCS = 138875005; descendants = 372405; IC = 1 - log(372405 + 1) / log(372406)
IC LCS = 0.000
JC = 1.000 + 0.411 - 2 * 0.000 = 1.411

S0.0 Concentration of neurotransmitter in blood [1285434004] -> C0.2 Body product observable [364684009]
concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC studente = 1.000
concetto corretto: descendants = 62; IC = 1 - log(62 + 1) / log(372406)
IC corretto = 0.677
LCS = 363787002; descendants = 10861; IC = 1 - log(10861 + 1) / log(372406)
IC LCS = 0.276
JC = 1.000 + 0.677 - 2 * 0.276 = 1.126

```

Figura 21. Calcolo dell'Information Content e della distanza Jiang-Conrath tra S0 e C0 Query2.

```

Matrice quadrata passata a Hungarian dentro SemanticEvaluator
      C0.0 Concentration of neur... [1285434004]  C0.1 Administrative statuses [307824009]  C0.2 Body prod
S0.0 Concentration of neur... [1285434004]  0.000  1.411  1.126
DUMMY-STUDENT  1.500  1.500  1.500
DUMMY-STUDENT  1.500  1.500  1.500
Assegnamento prodotto dentro SemanticEvaluator
S0.0 Concentration of neurotransmitter in blood [1285434004] -> C0.0 Concentration of neurotransmitter in blood [1285434004] | 0.000
DUMMY-STUDENT -> C0.1 Administrative statuses [307824009] | 1.500
DUMMY-STUDENT -> C0.2 Body product observable [364684009] | 1.500
Totale = 3.000

```

Figura 22. Matrice delle distanze e assegnamento interno tra S0 e C0 Query2.

Il parametro dello studente viene associato alla soluzione omonima con distanza pari a zero. Administrative statuses e Body product observable rimangono missing e contribuiscono con due penalità dummy. La somma interna è 3,000, divisa per i tre assegnamenti considerati, produce una distanza normalizzata tra S0 e C0 pari a 1,000.

Confronto tra S0 e C1 Query3

	C1.0 Carrier oil	C1.1 Oromucosal gel
S0.0 Neuro	LCS = 138875005 JC = 1.0 + 1.0 - 2*0.0 = 2.0	LCS = 138875005 JC = 1.0 + 0.946 - 2*0.0 = 1.946
Dummy	Dummy = 1.5	Dummy = 1.5

Tabella 25. Matrice delle distanze interne tra S0 e C1 Query3.

Assegnamento	Distanza
S0.0 Neurotransmitter → C1.1 Oromucosal gel	1.946
Dummy → C1.0 Carrier oil-containing product	1.5

Tabella 26. Assegnamento interno tra S0 e C1 Query3.

```

Primo livello: S0 contro C1 Query3

Calcolo delle distanze Jiang-Conrath

S0.0 Concentration of neurotransmitter in blood [1285434004] -> C1.0 Carrier oil-containing product [838458004]
concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC studente = 1.000
concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC corretto = 1.000
LCS = 138875005; descendants = 372405; IC = 1 - log(372405 + 1) / log(372406)
IC LCS = 0.000
JC = 1.000 + 1.000 - 2 * 0.000 = 2.000

S0.0 Concentration of neurotransmitter in blood [1285434004] -> C1.1 Oromucosal gel [385078007]
concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC studente = 1.000
concetto corretto: descendants = 1; IC = 1 - log(1 + 1) / log(372406)
IC corretto = 0.946
LCS = 138875005; descendants = 372405; IC = 1 - log(372405 + 1) / log(372406)
IC LCS = 0.000
JC = 1.000 + 0.946 - 2 * 0.000 = 1.946

```

Figura 23. Calcolo dell'Information Content e della distanza Jiang-Conrath tra S0 e C1 Query3.

Il concetto *Concentration of neurotransmitter in blood* viene associato a *Oromucosal gel* con distanza pari a 1,946, mentre *Carrier oil-containing product* rimane missing. La somma interna è 3,446, divisa per i due assegnamenti considerati, produce una distanza normalizzata tra S0 e C1 pari a 1,723.

La figura seguente mostra la matrice costruita per il confronto tra i parametri del gruppo S0 e quelli del gruppo corretto C1 Query3. Poiché S0 contiene un solo parametro, mentre C1 Query3 ne contiene due, viene aggiunta una riga dummy.

```

Matrice quadrata passata a Hungarian dentro SemanticEvaluator
                                C1.0 Carrier oil-containing... [838458004]  C1.1 Oromucosal gel [385078007]
S0.0 Concentration of neur... [1285434004]  2.000                        1.946
DUMMY-STUDENT                                1.500                        1.500
Assegnamento prodotto dentro SemanticEvaluator
S0.0 Concentration of neurotransmitter in blood [1285434004] -> C1.1 Oromucosal gel [385078007] | 1.946
DUMMY-STUDENT -> C1.0 Carrier oil-containing product [838458004] | 1.500
Totale = 3.446

```

Figura 24. Matrice delle distanze e assegnamento interno tra S0 e C1 Query3.

Assegnamento finale tra le data enquiry

Cella	Somma delle distanze e penalità	Assegnamenti considerati	Valore normalizzato
nodo S0 → nodo C0 Query2	3.000	3	1.000
nodo S0 → nodo C1 Query3	3.446	2	1.723

Tabella 27. Distanze normalizzate tra i gruppi di data enquiry del Segmento 2.

Gruppo studente	nodo C0 Query2	nodo C1 Query3
nodo S0	1.000	1.723
DUMMY-NODO-STUDENTE	1.500	1.500

Tabella 28. Matrice di assegnamento tra i gruppi delle data enquiry del Segmento 2.

Possibilità	Calcolo	Totale
A	nodo S0 → nodo C0 Query2 + DUMMY-NODO-STUDENTE → nodo C1 Query3 = 1.000 + 1.500	2.500
B	nodo S0 → nodo C1 Query3 + DUMMY-NODO-STUDENTE → nodo C0 Query2 = 1.723 + 1.500	3.223

Tabella 29. Confronto tra i possibili assegnamenti delle data enquiry del Segmento 2.

L'assegnamento tra i gruppi confronta le due possibili combinazioni. La soluzione S0-C0 con dummy-C1 ha un costo totale pari a 2,500, inferiore a 3,223 dell'alternativa. S0 viene quindi associato a C0 Query2, mentre C1 Query3 risulta missing.

	C0 Query2	C1 Query3
S0	1.000	1.723
DUMMY	1.500	1.500
CONFRONTO ASSEGNAMENTI		
A	S0 -> C0 Query2 + DUMMY -> C1 Query3 1.000 + 1.500 = 2.500 SCELTA	
B	S0 -> C1 Query3 + DUMMY -> C0 Query2 1.723 + 1.500 = 3.223	
RISULTATO MOSTRATO ALLO STUDENTE		
exact = 0 partial = 1 wrong = 0 missing = 1 extra = 0		
score = 0.167		

Figura 25. Matrice e assegnamento tra i gruppi delle data enquiry del Segmento 2.

Il gruppo S0-C0 è classificato come partial, poiché contiene un parametro exact e due parametri missing. Il suo score interno è quindi pari a $(1/3 = 0,333)$. Nel calcolo complessivo viene considerato anche il gruppo C1 Query3, che è missing e contribuisce con valore zero. Lo score delle data enquiry è pertanto $((0,333 + 0)/2 = 0,167)$.

5.3.3 Pharmacological prescription

La soluzione prevede due gruppi distinti: C0 contiene *Grain dust* e C1 contiene *Michel transport medium*. Lo studente inserisce entrambi i concetti nello stesso gruppo S0 e crea un secondo gruppo S1 contenente *Vegetable dust*, concetto semanticamente vicino a *Grain dust* ma non coincidente con esso.

Nessun gruppo viene associato in anticipo. Il confronto considera tutte le coppie S0-C0, S0-C1, S1-C0 e S1-C1. Per ciascuna coppia, la somma delle distanze interne e delle eventuali penalità viene normalizzata rispetto al numero degli assegnamenti interni. Le quattro distanze normalizzate formano la matrice utilizzata per individuare la combinazione con costo complessivo minimo.

Tipo	ID	Concetto	Path
Corretta C0	C0.0	Grain dust	Substance → Substance categorized by physical state → Dust → Organic dust → Vegetable dust → Grain dust
Corretta C1	C1.0	Michel transport medium	Substance → Substance categorized functionally → Specimen transport medium → Michel transport medium
Studente S0	S0.0	Grain dust	Substance → Substance categorized by physical state → Dust → Organic dust → Vegetable dust → Grain dust
Studente S0	S0.1	Michel transport medium	Substance → Substance categorized functionally → Specimen transport medium → Michel transport medium
Studente S1	S1.0	Vegetable dust	Substance → Substance categorized by physical state → Dust → Organic dust → Vegetable dust

Tabella 30. Concetti e path SNOMED CT dei gruppi di pharmacological prescription del Segmento 2.

Concetto	descendants	IC
304628004 Grain dust	0	$IC = 1 - \log(0 + 1) / \log(372406) = 1.0$
430028007 Michel transport medium	0	$IC = 1 - \log(0 + 1) / \log(372406) = 1.0$
304630002 Vegetable dust	17	$IC = 1 - \log(17 + 1) / \log(372406) = 0.775$
105590001 Substance	27675	$IC = 1 - \log(27675 + 1) / \log(372406) = 0.203$

Tabella 31. Valori di Information Content usati nei confronti delle pharmacological prescription.

Confronti del gruppo S0 con C0 e C1

Confronto tra gruppi	Parametri interni usati nel calcolo	Distanza
S0 - C0 Grain dust	S0.0 Grain → C0.0 Grain: LCS=304628004; JC = 1.0 + 1.0 - 2*1.0 = 0.0 S0.1 Michel → Dummy = 1.5	1.5
S0 - C1 Michel	S0.0 Grain → Dummy = 1.5 S0.1 Michel → C1.0 Michel: LCS=430028007; JC = 1.0 + 1.0 - 2*1.0 = 0.0	1.5

Tabella 32. Distanze interne tra S0 e i gruppi corretti C0 e C1.

```

Primo livello: S0 contro C0 Grain dust

Calcolo delle distanze Jiang-Conrath

S0.0 Grain dust [304628004] -> C0.0 Grain dust [304628004]
concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC studente = 1.000
concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC corretto = 1.000
LCS = 304628004; descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC LCS = 1.000
JC = 1.000 + 1.000 - 2 * 1.000 = 0.000

S0.1 Michel transport medium [430028007] -> C0.0 Grain dust [304628004]
concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC studente = 1.000
concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC corretto = 1.000
LCS = 105590001; descendants = 27675; IC = 1 - log(27675 + 1) / log(372406)
IC LCS = 0.203
JC = 1.000 + 1.000 - 2 * 0.203 = 1.595

```

Figura 26. Calcolo dell'Information Content e della distanza Jiang-Conrath tra S0 e C0 Grain dust.

```

S0.0 Grain dust [304628004] -> C1.0 Michel transport medium [430028007]
concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC studente = 1.000
concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC corretto = 1.000
LCS = 105590001; descendants = 27675; IC = 1 - log(27675 + 1) / log(372406)
IC LCS = 0.203
JC = 1.000 + 1.000 - 2 * 0.203 = 1.595

S0.1 Michel transport medium [430028007] -> C1.0 Michel transport medium [430028007]
concetto studente: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC studente = 1.000
concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC corretto = 1.000
LCS = 430028007; descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC LCS = 1.000
JC = 1.000 + 1.000 - 2 * 1.000 = 0.000

```

Figura 27. Calcolo dell'Information Content e della distanza Jiang-Conrath tra S0 e C1 Michel transport medium.

```

Matrice quadrata passata a Hungarian dentro SemanticEvaluator
                                C0.0 Grain dust [304628004]          DUMMY-CORRECT
S0.0 Grain dust [304628004]      0.000                            1.500
S0.1 Michel transport medium [430028007]  1.595                            1.500
Assegnamento prodotto dentro SemanticEvaluator
S0.0 Grain dust [304628004] -> C0.0 Grain dust [304628004] | 0.000
S0.1 Michel transport medium [430028007] -> DUMMY-CORRECT | 1.500
Totale = 1.500

```

```

Matrice quadrata passata a Hungarian dentro SemanticEvaluator
                                C1.0 Michel transport medium [430028007]  DUMMY-CORRECT
S0.0 Grain dust [304628004]      1.595                            1.500
S0.1 Michel transport medium [430028007]  0.000                            1.500
Assegnamento prodotto dentro SemanticEvaluator
S0.0 Grain dust [304628004] -> DUMMY-CORRECT | 1.500
S0.1 Michel transport medium [430028007] -> C1.0 Michel transport medium [430028007] | 0.000
Totale = 1.500

```

Figura 28. Confronti tra S0-C0 Grain dust e S0-C1 Michel transport medium.

Il gruppo S0 contiene due parametri, mentre ciascun gruppo corretto ne contiene uno. Nel confronto S0-C0, *Grain dust* è exact e *Michel transport medium* risulta extra, nel confronto S0-C1 avviene il contrario. In entrambi i casi la somma interna è 1,500 e, poiché gli assegnamenti considerati sono due, la distanza normalizzata tra i gruppi è pari a 0,750.

Confronti del gruppo S1 con C0 e C1

Confronto tra gruppi	Parametri interni usati nel calcolo	Distanza
S1-C0 Grain dust	S1.0 Vegetable dust → C0.0 Grain dust; JC = 0,225	0,225
S1-C1 Michel transport medium	S1.0 Vegetable dust → C1.0 Michel transport medium; JC = 1,369	1,369

Tabella 33. Distanze interne tra S1 e i gruppi corretti C0 e C1.

```

S1.0 Vegetable dust [304630002] -> C0.0 Grain dust [304628004]
concetto studente: descendants = 17; IC = 1 - log(17 + 1) / log(372406)
IC studente = 0.775
concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC corretto = 1.000
LCS = 304630002; descendants = 17; IC = 1 - log(17 + 1) / log(372406)
IC LCS = 0.775
JC = 0.775 + 1.000 - 2 * 0.775 = 0.225

```

Figura 29. Calcolo dell'Information Content e della distanza Jiang-Conrath tra S1 e C0 Grain dust.

```

S1.0 Vegetable dust [304630002] -> C1.0 Michel transport medium [430028007]
concetto studente: descendants = 17; IC = 1 - log(17 + 1) / log(372406)
IC studente = 0.775
concetto corretto: descendants = 0; IC = 1 - log(0 + 1) / log(372406)
IC corretto = 1.000
LCS = 105590001; descendants = 27675; IC = 1 - log(27675 + 1) / log(372406)
IC LCS = 0.203
JC = 0.775 + 1.000 - 2 * 0.203 = 1.369

```

Figura 30. Calcolo dell'Information Content e della distanza Jiang-Conrath tra S1 e C1 Michel transport medium.

```

Matrice quadrata passata a Hungarian dentro SemanticEvaluator
                                C0.0 Grain dust [304628004]
S1.0 Vegetable dust [304630002] 0.225
Assegnamento prodotto dentro SemanticEvaluator
S1.0 Vegetable dust [304630002] -> C0.0 Grain dust [304628004] | 0.225
Totale = 0.225

Primo livello: S1 contro C1 Michel transport medium

Matrice JC creata dentro SemanticEvaluator
                                C1.0 Michel transport medium [430028007]
S1.0 Vegetable dust [304630002] 1.369

Matrice quadrata passata a Hungarian dentro SemanticEvaluator
                                C1.0 Michel transport medium [430028007]
S1.0 Vegetable dust [304630002] 1.369
Assegnamento prodotto dentro SemanticEvaluator
S1.0 Vegetable dust [304630002] -> C1.0 Michel transport medium [430028007] | 1.369
Totale = 1.369

```

Figura 31. Confronti tra S1-C0 Grain dust e S1-C1 Michel transport medium.

Il concetto *Vegetable dust* produce una distanza pari a 0,225 rispetto a *Grain dust* e pari a 1,369 rispetto a *Michel transport medium*. Poiché ciascun confronto contiene un solo assegnamento interno, i valori normalizzati rimangono rispettivamente 0,225 e 1,369.

Assegnamento finale tra i gruppi

Cella	Somma delle distanze e penalità	Assegnamenti considerati	Valore normalizzato
nodo S0 → nodo C0 Grain dust	1.500	2	0.750
nodo S0 → nodo C1 Michel transport medium	1.500	2	0.750
nodo S1 → nodo C0 Grain dust	0.225	1	0.225
nodo S1 → nodo C1 Michel transport medium	1.369	1	1.369

Tabella 34. Distanze normalizzate delle quattro coppie di pharmacological prescription.

Gruppo studente	nodo C0 Grain dust	nodo C1 Michel transport medium
nodo S0	0.750	0.750
nodo S1	0.225	1.369

Tabella 35. Matrice di assegnamento tra i gruppi delle pharmacological prescription del Segmento 2.

Possibilità	Calcolo	Totale
A	nodo S0 → nodo C0 + nodo S1 → nodo C1 = 0.750 + 1.369	2.119
B	nodo S0 → nodo C1 + nodo S1 → nodo C0 = 0.750 + 0.225	0.975

Tabella 36. Confronto tra i possibili assegnamenti delle pharmacological prescription del Segmento 2.

L'assegnamento S0-C1 e S1-C0 produce un costo pari a 0,975, inferiore al valore 2,119 dell'altra combinazione. L'assegnamento con costo minimo associa pertanto S0 a *Michel transport medium* e S1 a *Grain dust*.

	C0 Grain dust	C1 Michel transport medium
S0	0.750	0.750
S1	0.225	1.369
CONFRONTO ASSEGNAMENTI		
A	S0 -> C0 Grain dust + S1 -> C1 Michel transport medium 0.750 + 1.369 = 2.119	
B	S0 -> C1 Michel transport medium + S1 -> C0 Grain dust 0.750 + 0.225 = 0.975 SCELTA	
RISULTATO MOSTRATO ALLO STUDENTE		
exact = 0 partial = 1 wrong = 1 missing = 0 extra = 0		
score = 0.250		

Figura 32. Matrice e assegnamento tra i gruppi delle pharmacological prescription del Segmento 2.

Nel gruppo S0, *Michel transport medium* è exact e *Grain dust* è extra interno, il gruppo viene classificato come partial e contribuisce con lo score interno 0,5. Il gruppo S1 non contiene parametri exact e viene classificato come wrong. Lo score complessivo delle pharmacological prescription è 0,250.

5.4 Flusso della verifica dal punto di vista dello studente

Questa sezione descrive il funzionamento della verifica dal punto di vista dello studente, seguendo l'ordine delle operazioni mostrate nell'interfaccia: avvio dell'esercizio, inserimento delle risposte, visualizzazione del risultato e aggiornamento progressivo del pathway.

5.4.1 Avvio della verifica

All'avvio dell'esercizio l'applicazione carica il caso clinico e la guideline associata. Il percorso corretto è disponibile al sistema, ma non viene mostrato allo studente prima della risposta.

Il pathway viene suddiviso in segmenti delimitati dall'inizio del percorso o da una decisione precedente e dalla decisione successiva. Sono sottoposti alla verifica i segmenti che contengono almeno una work action, una data enquiry o una pharmacological prescription.

La parte corretta del grafo viene rivelata progressivamente dopo la valutazione dei segmenti. Il pathway completo diventa visibile al termine dell'esercizio, quando non può più anticipare le risposte successive.

5.4.2 Svolgimento della verifica nel Segmento 1

Inserimento delle risposte

Nel primo segmento lo studente inserisce le risposte ritenute corrette utilizzando le sezioni dedicate ai diversi tipi di nodo supportati.

Figura 33. Segmento 1: schermata di verifica con le sezioni in cui lo studente inserisce le risposte.

Le work action corrispondono a singoli concetti SNOMED CT. Le data enquiry e le pharmacological prescription sono invece gruppi composti, all'interno dei quali lo studente può aggiungere uno o più parametri. Ogni concetto selezionato, sia come work action sia come parametro di un nodo composto, viene visualizzato nella relativa sezione del pannello insieme al path SNOMED CT che conduce al concetto scelto.

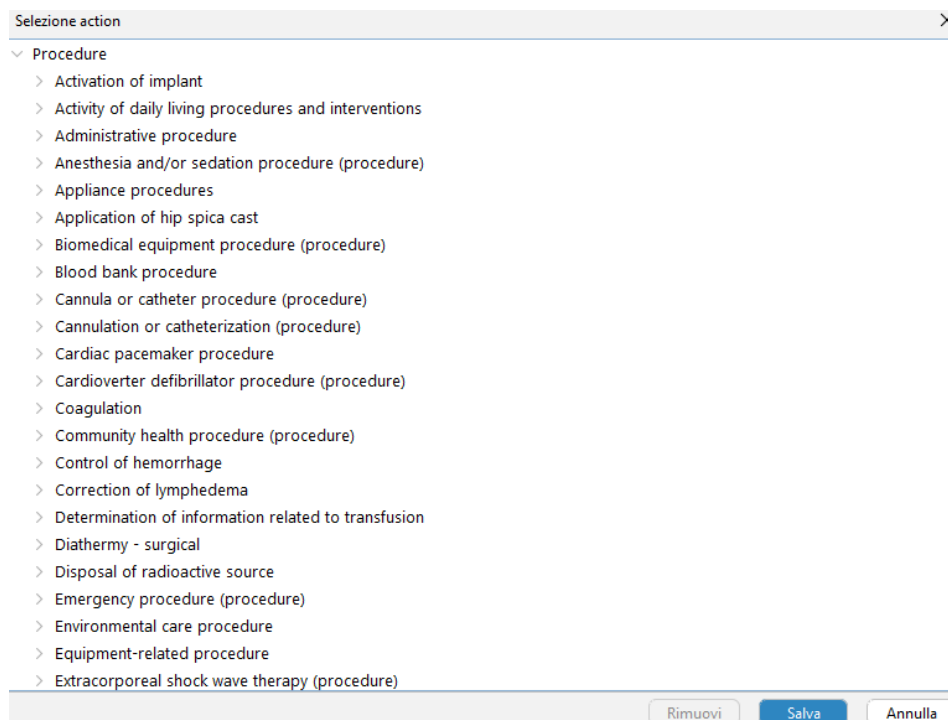


Figura 34. Selezione di una work action nell'albero SNOMED.

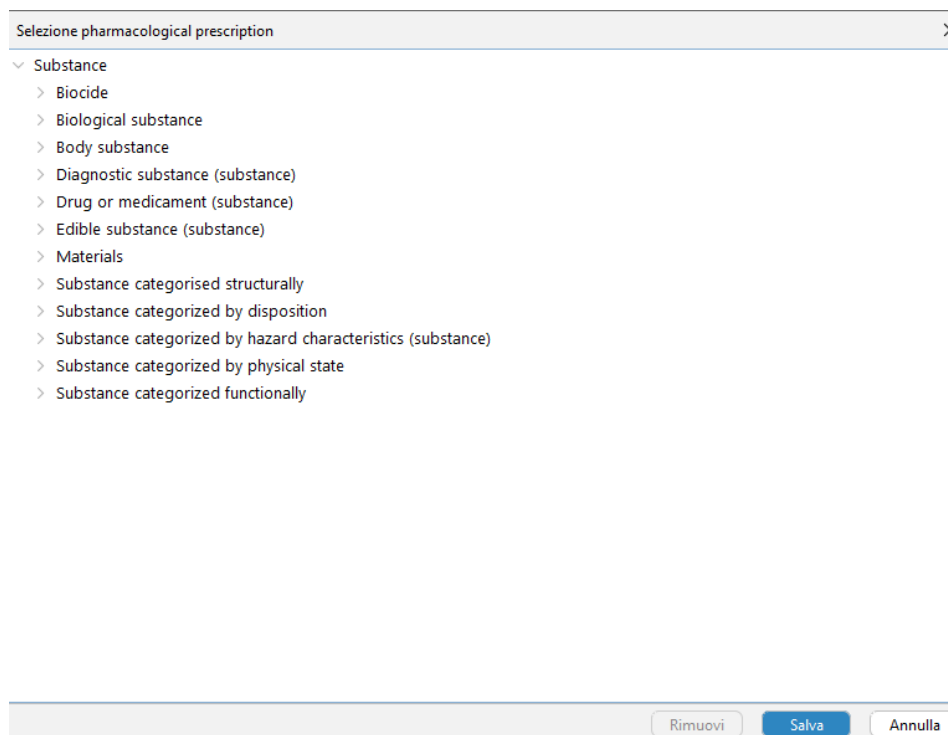


Figura 35. Selezione di una pharmacological prescription nell'albero SNOMED.

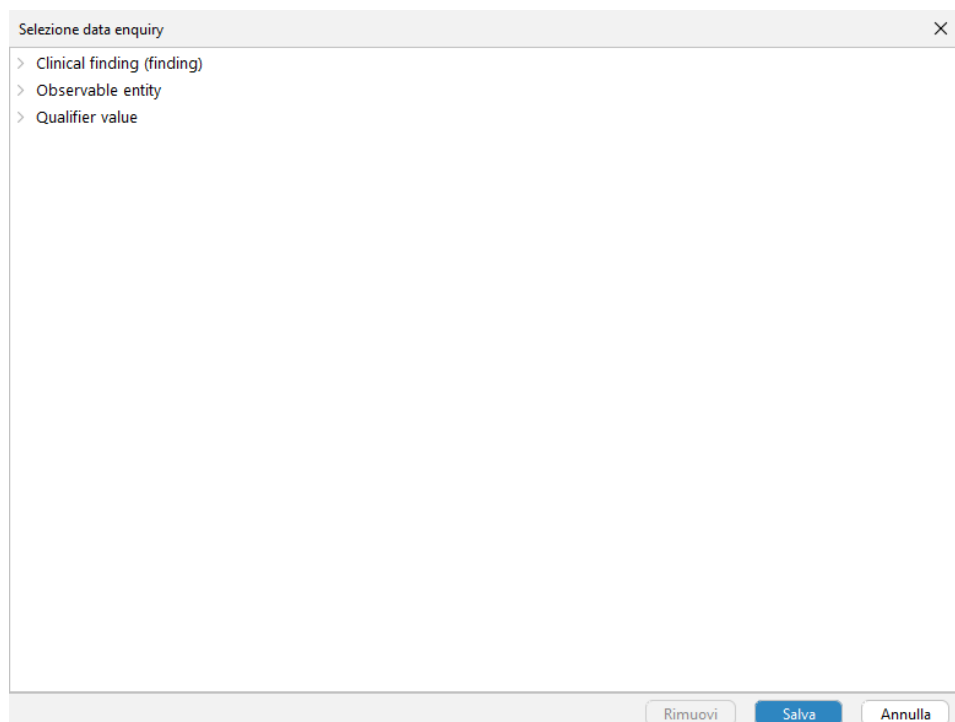


Figura 36. Selezione di una data enquiry nell'albero SNOMED.

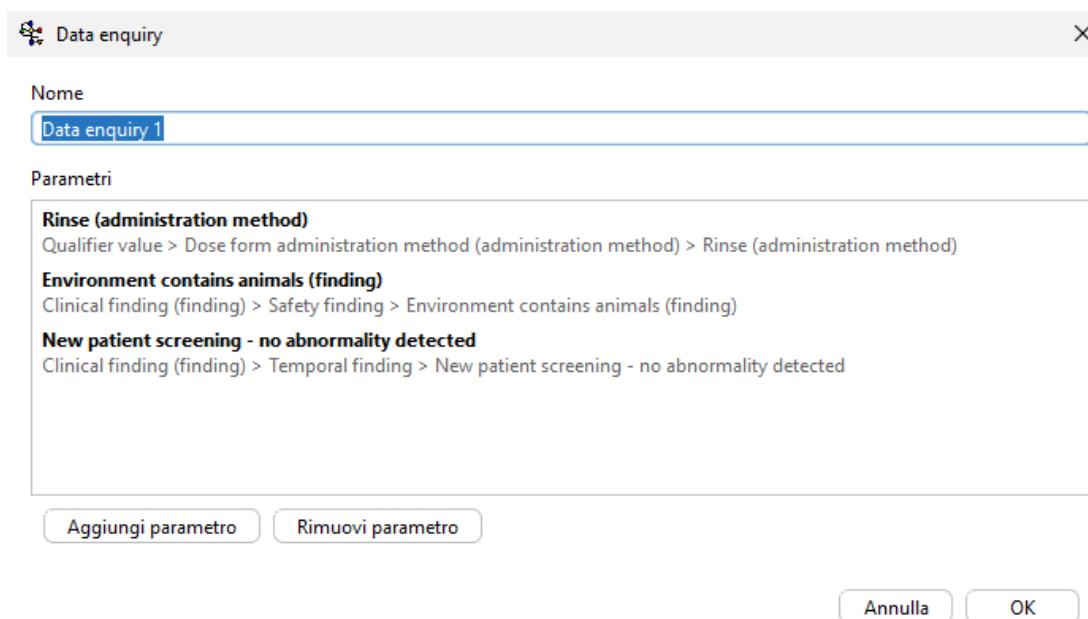


Figura 37. Segmento 1: composizione del nodo Data enquiry 1 con tre parametri interni.

Data enquiry [X]

Nome
Data enquiry 2

Parametri

Asymptomatic
Clinical finding (finding) > Asymptomatic

Aggiungi parametro Rimuovi parametro

Annulla OK

Figura 38. Segmento 1: composizione del nodo Data enquiry 2 con un parametro interno.

Risultato e dettaglio delle associazioni

Dopo l'invio, i controlli di modifica vengono disabilitati e il segmento viene valutato. Il riepilogo mostra i conteggi exact, partial, wrong, missing ed extra, insieme allo score ottenuto.

Il riepilogo mostra l'esito generale della verifica attraverso i conteggi e lo score. Il dettaglio consente poi di esaminare le associazioni individuate dal sistema e la classificazione assegnata a ciascuna risposta. Per i nodi composti vengono mostrati anche i risultati dei parametri interni, dai quali deriva la classificazione complessiva del gruppo.

Risultato verifica segmento 1

	Action	Data enquiry	Pharmacological prescription
✓ Exact	2	0	0
⚙ Partial	0	1	0
✗ Wrong	0	1	0
! Missing	0	0	0
! Extra	1	0	0
📊 Score	67%	17%	0%

Figura 39. Segmento 1: riepilogo del risultato con conteggi e score.

Sotto il riepilogo, la sezione *Soluzioni corrette* elenca gli elementi previsti nel segmento. Per ogni work action e per ogni parametro delle data enquiry e delle pharmacological prescription vengono mostrati il nome del concetto e, sotto di esso, il relativo path SNOMED CT fino al concetto corretto.

Soluzioni corrette	
Action	
* Haematopoietic stem cell mobilisation	Procedure > Stem cell mobilization > Haematopoietic stem cell mobilisation
* Coagulation	Procedure > Coagulation
Data enquiry	
Data enquiry 1	
* Rinse (administration method)	Qualifier value > Dose form administration method (administration method) > Rinse (administration method)
* No consent for MR - Measles/rubella vaccine	Clinical finding (finding) > Temporal finding > No consent for MR - Measles/rubella vaccine
Data enquiry 2	
* Clinical privilege (finding)	Clinical finding (finding) > Administrative statuses > Clinical privilege (finding)
* Retigabine	Substance > Drug or medicament (substance) > Central nervous system agent (substance) > Anticonvulsant > Retigabine
* Identification of fungus by MALDI-TOF MS (matrix assisted laser desorption ionization time of flight mass spectrometry)	Observable entity > Identification of fungus by MALDI-TOF MS (matrix assisted laser desorption ionization time of flight mass spectrometry)
Pharmacological prescription	

Figura 40. Segmento 1: soluzioni corrette mostrate allo studente.

La sezione di dettaglio presenta le associazioni selezionate dal matching. Nei nodi composti sono indicati anche i parametri interni exact, wrong, missing ed extra, così da rendere comprensibile la classificazione del gruppo.

Dettaglio	
✓ Exact (2)	
Action	
Tua risposta: Haematopoietic stem cell mobilisation	
Corretta: Haematopoietic stem cell mobilisation	
Action	
Tua risposta: Coagulation	
Corretta: Coagulation	
📑 Partial (1)	
Data enquiry	
Tua risposta: Data enquiry 1: Rinse (administration method), Environment contains animals (finding), New patient screening - no abnormality detected	
Corretta: Corretta 1: Rinse (administration method), No consent for MR - Measles/rubella vaccine	
Exact - tua risposta: Rinse (administration method) corretta: Rinse (administration method)	
Wrong - tua risposta: New patient screening - no abnormality detected corretta: No consent for MR - Measles/rubella vaccine	
Extra - Environment contains animals (finding)	
Dettaglio SNOMED	
✗ Wrong (1)	
Data enquiry	
Tua risposta: Data enquiry 2: Asymptomatic	
Corretta: Corretta 2: Clinical privilege (finding), Retigabine, Identification of fungus by MALDI-TOF MS (matrix assisted laser desorption ionization time of flight mass spectrometry)	
Wrong - tua risposta: Asymptomatic corretta: Clinical privilege (finding)	
Missing - Retigabine	
Missing - Identification of fungus by MALDI-TOF MS (matrix assisted laser desorption ionization time of flight mass spectrometry)	
Dettaglio SNOMED	
! Extra (1)	
Action	
Tua risposta: Medical therapy	

Figura 41. Segmento 1: dettaglio della valutazione con exact, partial, wrong ed extra.

Dettaglio SNOMED e aggiornamento del grafo

Per le risposte wrong e per i gruppi partial è possibile aprire il Dettaglio SNOMED, che affianca il path selezionato dallo studente e quello corretto. La visualizzazione permette di riconoscere gli antenati condivisi e il punto in cui i due percorsi divergono.

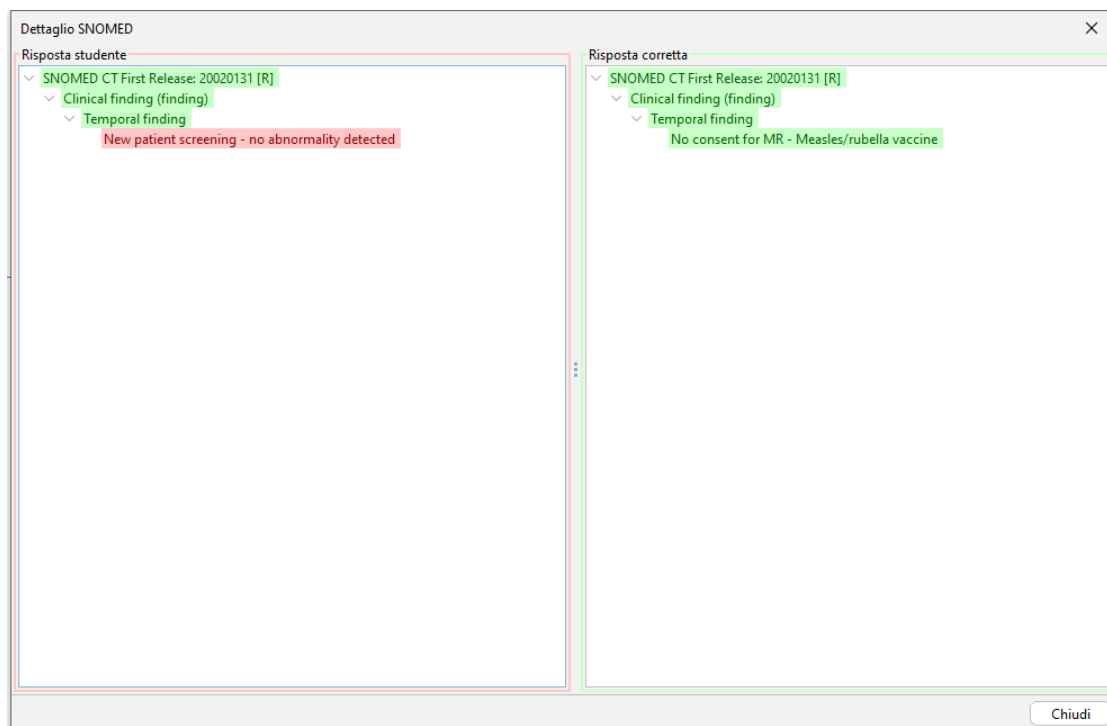


Figura 42. Segmento 1: Dettaglio SNOMED tra *New patient screening - no abnormality detected* e *No consent for MR*.

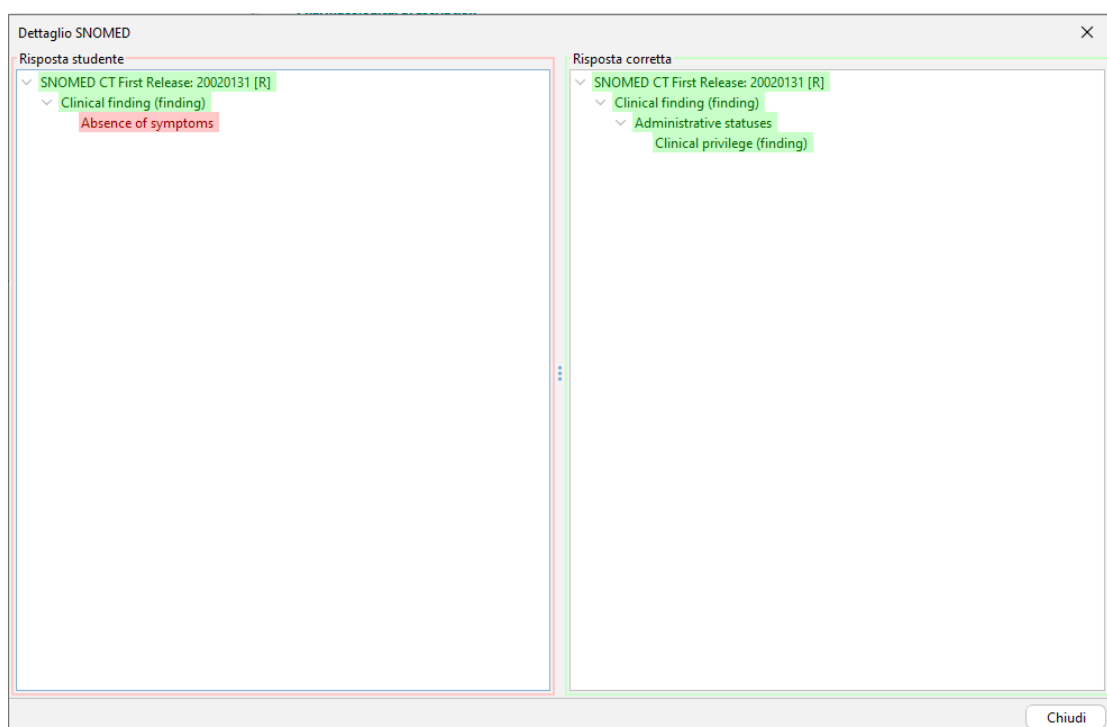


Figura 43. Segmento 1: Dettaglio SNOMED tra *Asymptomatic* e *Clinical privilege*.

Dopo la valutazione, il grafo viene aggiornato mostrando i nodi e gli archi relativi al segmento appena verificato. In questo modo, lo studente può individuare nel pathway la posizione degli elementi a cui si riferisce il feedback.

La visualizzazione del grafo non sostituisce il feedback semantico, ma lo completa. Il dettaglio SNOMED spiega perché due concetti sono stati associati e dove divergono nella gerarchia, il pathway mostra invece il ruolo degli elementi corretti all'interno della sequenza clinica. I due livelli vengono presentati solo dopo la valutazione del segmento, evitando di anticipare le soluzioni successive.

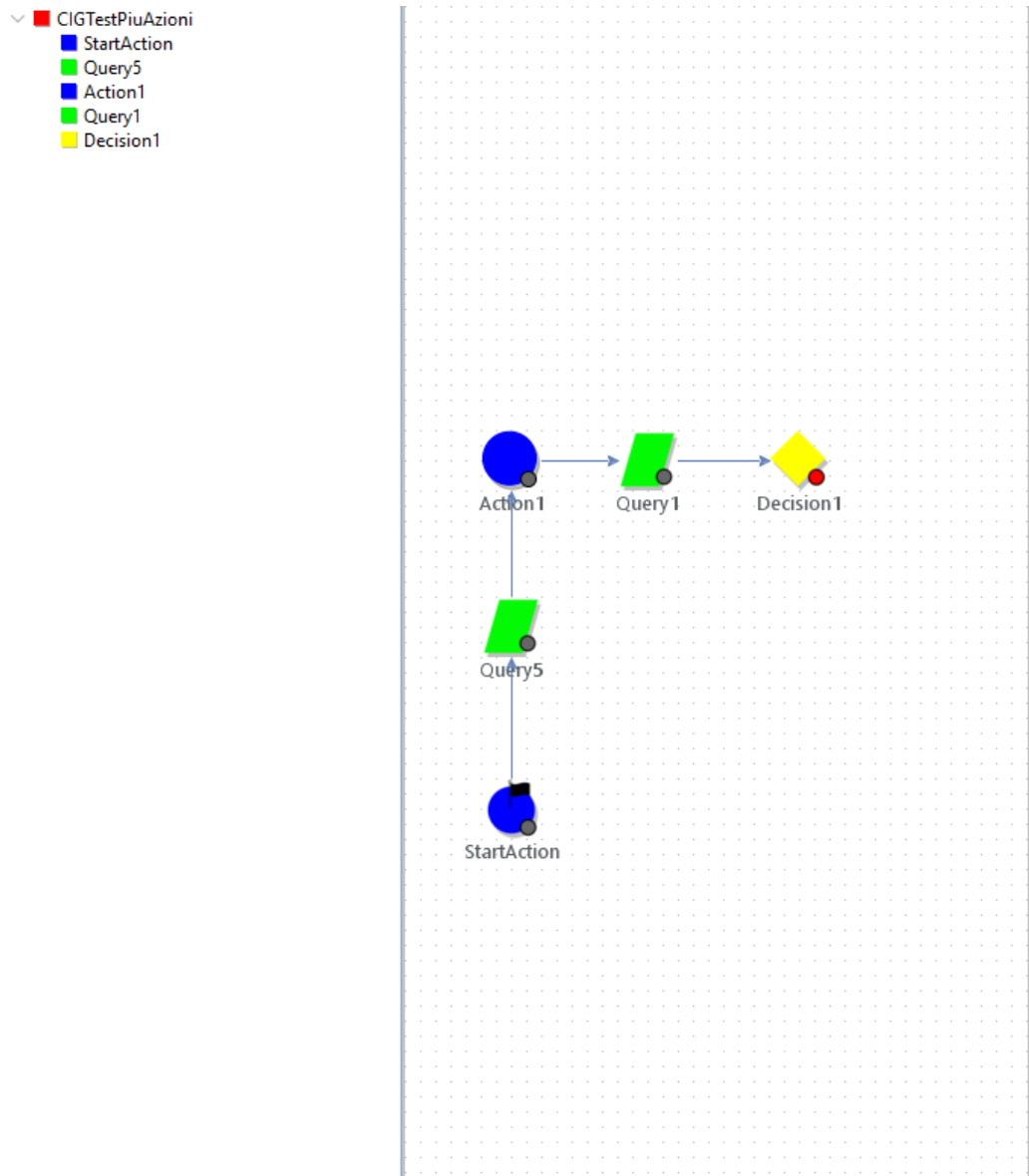


Figura 44. Segmento 1: porzione di grafo rivelata dopo la valutazione del segmento.

5.4.3 Svolgimento della verifica nel Segmento 2

Inserimento delle risposte

Nel secondo segmento lo studente utilizza nuovamente le tre sezioni del pannello. La parte del percorso già valutata rimane visibile, mentre i nodi corretti del segmento corrente restano nascosti fino all'invio delle risposte.

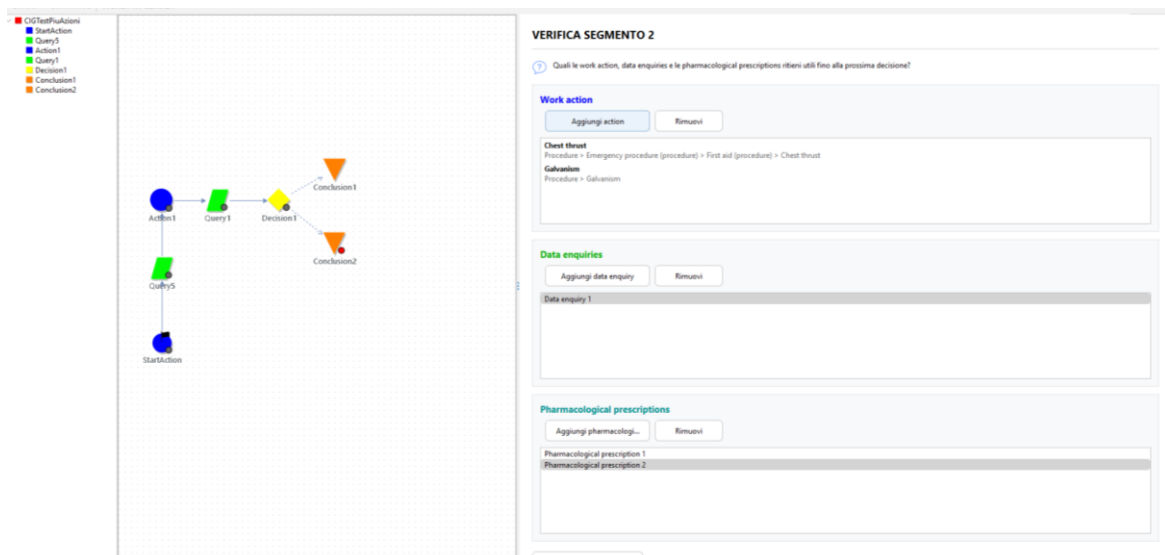


Figura 45. Segmento 2: schermata di verifica con grafo parziale a sinistra e pannello di risposta a destra.

Nell'esempio viene inserita una sola data enquiry, composta da un unico parametro interno.

Figura 46. Segmento 2: composizione della data enquiry inserita dallo studente.

Per le *pharmacological prescription* la soluzione prevede due gruppi distinti: il primo contiene *Grain dust*, mentre il secondo contiene *Michel transport medium*. Lo studente inserisce invece entrambi i concetti nello stesso gruppo e crea un secondo gruppo contenente *Vegetable dust*, concetto semanticamente vicino a *Grain dust* ma non coincidente con esso. L'esempio mostra che la presenza di concetti corretti non è sufficiente se questi vengono organizzati in gruppi diversi da quelli previsti dalla soluzione.

Pharmacological prescription

Nome
Pharmacological prescription 1

Parametri

Grain dust
Substance > Substance categorized by physical state > Dust > Organic dust > Vegetable dust > Grain dust

Michel transport medium (substance)
Substance > Substance categorized functionally > Specimen transport medium > Michel transport medium (substance)

Aggiungi parametro Rimuovi parametro

Annulla OK

Figura 47. Segmento 2: Pharmacological prescription 1 con Grain dust e Michel transport medium nello stesso gruppo.

Pharmacological prescription

Nome
Pharmacological prescription 2

Parametri

Vegetable dust
Substance > Substance categorized by physical state > Dust > Organic dust > Vegetable dust

Aggiungi parametro Rimuovi parametro

Annulla OK

Figura 48. Segmento 2: Pharmacological prescription 2 con Vegetable dust.

Risultato e dettaglio delle associazioni

Il riepilogo del secondo segmento comprende risultati partial, wrong e missing. Alcune risposte sono semanticamente vicine alle soluzioni ma non esatte, altri elementi corretti non sono stati inseriti.

Risultato verifica segmento 2

	Action	Data enquiry	Pharmacological prescription
✓ Exact	0	0	0
⚙ Partial	0	1	1
✗ Wrong	2	0	1
! Missing	1	1	0
! Extra	0	0	0
📊 Score	0%	17%	25%

Figura 49. Segmento 2: riepilogo del risultato con conteggi e score.

La sezione *Soluzioni corrette* elenca le work action, le data enquiry e le pharmacological prescription previste nel segmento.

Soluzioni corrette	
Action	
• Emergency procedure (procedure)	Procedure > Emergency procedure (procedure)
• Dysphagia therapy regime	Procedure > Regimes and therapies > Dysphagia therapy regime
• Application of hip spica cast	Procedure > Application of hip spica cast
Data enquiry	
Data enquiry 1	
• Concentration of neurotransmitter in blood	Observable entity > Concentration of neurotransmitter in blood
• Administrative statuses	Clinical finding (finding) > Administrative statuses
• Body product observable	Observable entity > Body product observable
Data enquiry 2	
• Carrier oil-containing product	Qualifier value > Role > Carrier oil > Carrier oil-containing product
• Oromucosal gel	Qualifier value > Pharmaceutical dose form (dose form) > Oromucosal dose form (dose form) > Oromucosal gel
Pharmacological prescription	
Pharmacological prescription 1	
• Grain dust	Substance > Substance categorized by physical state > Dust > Organic dust > Vegetable dust > Grain dust
Pharmacological prescription 2	
• Michel transport medium (substance)	Substance > Substance categorized functionally > Specimen transport medium > Michel transport medium (substance)

Figura 50. Segmento 2: soluzioni corrette mostrate allo studente.

Nel dettaglio vengono mostrate le coppie selezionate dall'algoritmo, tra cui *Chest thrust* con *Emergency procedure*, *Galvanism* con *Dysphagia therapy regime* e *Vegetable dust* con *Grain dust*.

Dettaglio

Partial (2)

Data enquiry
 Tua risposta: **Data enquiry 1: Concentration of neurotransmitter in blood**
 Corretta: **Corretta 1: Concentration of neurotransmitter in blood, Administrative statuses, Body product observable**
 Exact - tua risposta: Concentration of neurotransmitter in blood | corretta: Concentration of neurotransmitter in blood
 Missing - Administrative statuses
 Missing - Body product observable

Pharmacological prescription
 Tua risposta: **Pharmacological prescription 1: Grain dust, Michel transport medium (substance)**
 Corretta: **Corretta 2: Michel transport medium (substance)**
 Exact - tua risposta: Michel transport medium (substance) | corretta: Michel transport medium (substance)
 Extra - Grain dust

Wrong (3)

Action
 Tua risposta: **Chest thrust**
 Corretta: **Emergency procedure (procedure)**
 Dettaglio SNOMED

Action
 Tua risposta: **Galvanism**
 Corretta: **Dysphagia therapy regime**
 Dettaglio SNOMED

Pharmacological prescription
 Tua risposta: **Pharmacological prescription 2: Vegetable dust**
 Corretta: **Corretta 1: Grain dust**
 Wrong - tua risposta: Vegetable dust | corretta: Grain dust
 Dettaglio SNOMED

Missing (2)

Action
 Corretta: **Application of hip spica cast**

Data enquiry
 Corretta: **Corretta 2: Carrier oil-containing product, Oromucosal gel**

Figura 51. Segmento 2: dettaglio della valutazione con partial, wrong e missing.

Dettaglio SNOMED e aggiornamento del grafo

Il *Dettaglio SNOMED* è disponibile anche per le associazioni del secondo segmento. La figura seguente mostra il confronto tra il concetto inserito nella seconda pharmacological prescription e quello corretto associato dal sistema.

Dettaglio SNOMED

Risposta studente

- SNOMED CT First Release: 20020131 [R]
 - Substance
 - Substance categorized by physical state
 - Dust
 - Organic dust
 - Vegetable dust

Risposta corretta

- SNOMED CT First Release: 20020131 [R]
 - Substance
 - Substance categorized by physical state
 - Dust
 - Organic dust
 - Vegetable dust
 - Grain dust

Chiudi

Figura 52. Segmento 2: Dettaglio SNOMED tra Vegetable dust e Grain dust.

Al termine della valutazione viene rivelata anche la porzione del pathway relativa al secondo segmento.

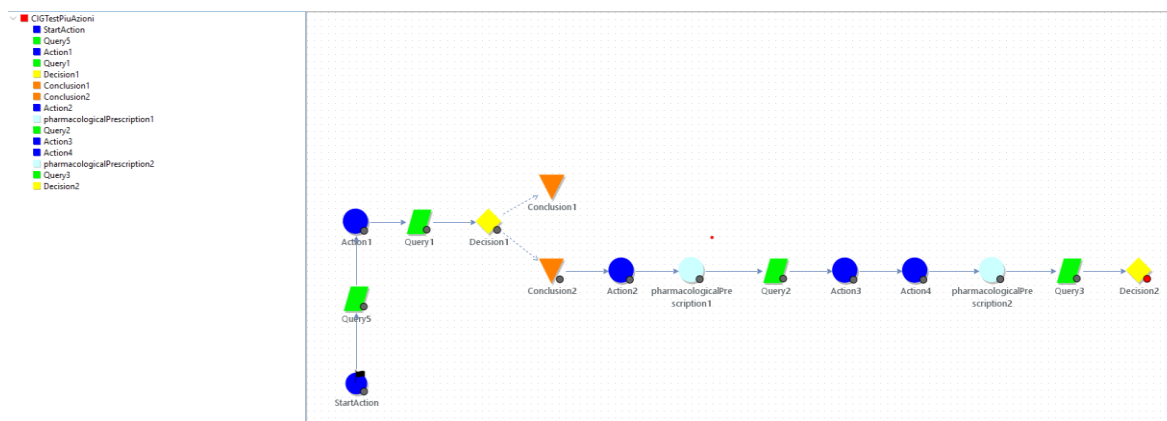


Figura 53. Segmento 2: porzione di grafo rivelata dopo la valutazione.

5.4.4 Conclusione dell'esercizio e riepilogo

Dopo il completamento di tutti i segmenti, l'applicazione mostra il pathway completo. La visualizzazione assume a questo punto una funzione di riepilogo e permette allo studente di ricostruire l'intero percorso corretto.

La schermata conclusiva riunisce le informazioni che durante l'esercitazione erano state mostrate in modo progressivo. Il pathway completo fornisce la visione strutturale della guideline, mentre il riepilogo conserva il dettaglio delle decisioni di matching e delle classificazioni prodotte in ciascun segmento.

Il comando *Show Summary* consente di rivedere i segmenti già svolti, le risposte inserite, le soluzioni corrette, le classificazioni, gli score e i confronti disponibili nel *Dettaglio SNOMED*.

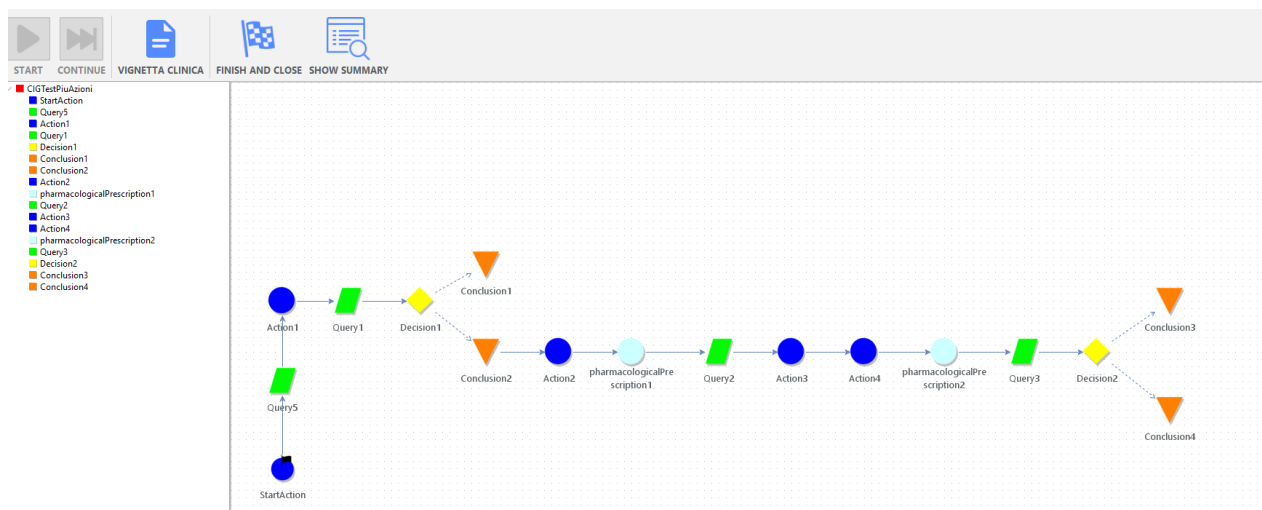


Figura 54. Fine dell'esercizio: guideline completa visibile e pulsante *Show Summary* per consultare il riepilogo dei segmenti svolti.

Gli esempi mostrano anche il ruolo degli elementi dummy. Questi completano la matrice quando il numero delle risposte dello studente è diverso dal numero degli elementi corretti. Le risposte che non trovano corrispondenza vengono classificate come extra, mentre gli elementi corretti non coperti vengono classificati come missing. Le relative penalità vengono considerate sia nel matching sia nel calcolo dello score. Nei nodi composti lo stesso procedimento viene applicato prima ai parametri interni e successivamente ai gruppi.

I valori riportati nelle tabelle coincidono con i risultati visualizzati dall'applicazione. Oltre allo score complessivo, lo studente può consultare le associazioni prodotte dal matching, gli elementi mancanti o aggiuntivi e il confronto tra i path SNOMED CT. Il feedback permette così di comprendere da quali elementi deriva la valutazione ricevuta.

6 Discussione e conclusioni

Questo capitolo conclude il lavoro discutendo il possibile impiego didattico del modulo, gli aspetti da considerare nella sua applicazione e le possibili direzioni di sviluppo del sistema. Particolare attenzione è dedicata al ruolo del feedback nel processo di valutazione formativa e alle informazioni che può offrire allo studente.

6.1 Valore didattico del feedback

Il feedback formativo può essere definito come un'informazione fornita allo studente con lo scopo di modificare il suo modo di pensare o di agire e favorire l'apprendimento. In questo ambito si distingue tra *verification*, che comunica la correttezza della risposta, ed *elaboration*, che aggiunge informazioni utili per comprenderla e rivederla. Il feedback non dovrebbe pertanto limitarsi a indicare l'esito della verifica, ma offrire elementi riferiti alla risposta e al compito svolto¹⁹.

Nel modulo sviluppato, le categorie di valutazione e lo score forniscono l'esito della verifica, mentre le associazioni, il dettaglio dei nodi composti, il confronto tra i path SNOMED CT e la visualizzazione del segmento permettono di approfondirne il significato. Lo studente può così individuare la parte della risposta che si discosta dalla soluzione attesa e osservare la relazione tra i concetti confrontati.

Un altro aspetto rilevante è la specificità delle informazioni fornite. Le informazioni riferite alla risposta fornita possono essere più utili di indicazioni generiche, perché rendono più chiaro ciò che deve essere corretto o approfondito.¹⁹ Nel modulo, il risultato rimane collegato agli elementi inseriti nello specifico segmento e, nel caso dei nodi composti, ai rispettivi parametri interni.

La quantità di informazioni deve però rimanere adeguata al compito. Un messaggio troppo lungo o complesso può rendere meno chiaro il contenuto principale e ridurne l'utilità¹⁹. La presenza di un riepilogo iniziale e di dettagli consultabili separatamente permette di mantenere visibile l'esito generale, lasciando allo studente la possibilità di approfondire le associazioni e i path quando necessario.

Nel complesso, il modulo non si limita a comunicare l'esito della verifica, ma fornisce informazioni collegate alla risposta dello studente e consultabili a diversi livelli di dettaglio.

6.2 Aspetti da considerare nell'applicazione del modulo

Il funzionamento del modulo è stato verificato attraverso i due segmenti analizzati nel Capitolo 5, costruiti per comprendere le principali situazioni gestite dalla valutazione. L'applicazione ad altri case study e a guideline con strutture differenti permetterebbe di osservare il comportamento del metodo in scenari più vari e di individuare eventuali casi limite.

Un altro aspetto da considerare riguarda la versione di SNOMED CT utilizzata dal sistema. Le edizioni di SNOMED CT vengono aggiornate regolarmente e ogni nuova versione può aggiungere, modificare o rendere inattivi concetti, relazioni e membri dei reference set. Di conseguenza, i risultati delle operazioni terminologiche possono variare in base alla versione selezionata.²⁰

Per garantire la riproducibilità dei risultati, è opportuno indicare la versione della terminologia e della transitive closure impiegate durante la valutazione e verificare la compatibilità dei riferimenti ontologici in caso di aggiornamento.

6.3 Validazione e sviluppi futuri

Una fase successiva del lavoro potrebbe riguardare l'utilizzo del modulo da parte di studenti e docenti. L'osservazione delle esercitazioni permetterebbe di verificare se le categorie di valutazione, lo score, le associazioni e il Dettaglio SNOMED risultano comprensibili. Le osservazioni raccolte potrebbero essere utilizzate per migliorare la presentazione del feedback, evidenziando le informazioni considerate più utili e semplificando quelle che risultano poco chiare o troppo dettagliate.

Un possibile sviluppo riguarda la modalità con cui lo studente costruisce la propria risposta. Attualmente vengono selezionati i concetti SNOMED CT ritenuti corretti per il segmento. In futuro, lo studente potrebbe invece ricostruire direttamente una propria versione del segmento, inserendo i nodi e gli archi che rappresentano il percorso clinico.

Il confronto con la guideline non riguarderebbe più soltanto la presenza dei concetti, ma anche la loro organizzazione all'interno del segmento. Il sistema potrebbe valutare l'ordine delle attività, i nodi mancanti o aggiuntivi, gli archi inseriti e, quando presenti, le relazioni temporali tra le azioni. La verifica continuerebbe a essere svolta segmento per segmento, ma prenderebbe in considerazione anche la struttura del percorso costruito dallo studente.

Questa estensione renderebbe l'esercitazione più vicina alla simulazione e alla costruzione di una guideline. Sarebbe però necessario definire nuove modalità di confronto tra il segmento costruito dallo studente e quello previsto dal case study, mantenendo separata la valutazione semantica dei concetti dalla valutazione della struttura composta da nodi e archi.

Un ulteriore sviluppo riguarda le pharmacological prescription. Nel modulo attuale la valutazione considera i concetti SNOMED CT inseriti nei gruppi. In una versione futura, ogni prescrizione potrebbe includere informazioni più dettagliate, come la dose o la concentrazione del farmaco, la frequenza e la via di somministrazione e la durata del trattamento. Il sistema potrebbe distinguere la correttezza del farmaco selezionato dalla correttezza dei parametri che ne descrivono l'utilizzo.

6.4 Considerazioni conclusive

Il presente lavoro ha affrontato il problema della valutazione delle risposte dello studente in GLARE-Edu, introducendo un confronto che non dipende esclusivamente dalla coincidenza tra gli elementi inseriti e quelli previsti dalla guideline. A questo scopo è stato progettato e integrato un modulo che utilizza la struttura gerarchica di SNOMED CT per calcolare la distanza semantica e individuare le corrispondenze tra le risposte dello studente e gli elementi attesi.

Il contributo del lavoro non riguarda soltanto la scelta di una misura semantica, ma la sua applicazione all'intero flusso di verifica. Il metodo è stato adattato alla suddivisione del pathway in segmenti e alle diverse strutture dei nodi considerati, fino alla realizzazione di una componente operativa utilizzabile nell'interfaccia di GLARE-Edu.

Gli esempi analizzati mostrano che il modulo consente di distinguere corrispondenze esatte, differenze semantiche, elementi mancanti ed elementi aggiuntivi, mantenendo il risultato collegato alla porzione di pathway valutata. Questi risultati costituiscono una verifica funzionale dell'approccio, ma non dimostrano ancora la sua efficacia didattica. Il passo successivo consiste in una sperimentazione con studenti e docenti, finalizzata a valutare la comprensibilità del feedback, l'adeguatezza dello score e l'utilità delle informazioni semantiche durante l'esercitazione.

Riferimenti bibliografici

1. Wasylewicz, A. T. M. & Scheepers-Hoeks, A. M. J. W. Clinical Decision Support Systems. in *Fundamentals of Clinical Data Science* (eds Kubben, P., Dumontier, M. & Dekker, A.) (Springer, Cham (CH), 2019).
2. De Clercq, P., Kaiser, K. & Hasman, A. Computer-Interpretable Guideline formalisms. *Stud. Health Technol. Inform.* **139**, 22–43 (2008).
3. Peleg, M. Computer-interpretable clinical guidelines: a methodological review. *J. Biomed. Inform.* **46**, 744–763 (2013).
4. Bottrighi, A. *et al.* A Symbolic AI Approach to Medical Training. *J. Med. Syst.* **49**, 2 (2025).
5. SNOMED CT Basics | Practical Guides | SNOMED International Documents. <https://docs.snomed.org/snomed-ct-practical-guides/snomed-ct-starter-guide/4-snomed-ct-basics> (2025).
6. SNOMED CT Logical Model | Practical Guides | SNOMED International Documents. <https://docs.snomed.org/snomed-ct-practical-guides/snomed-ct-starter-guide/5-snomed-ct-logical-model> (2025).
7. Nicol, D. J. & Macfarlane-Dick, D. Formative assessment and self-regulated learning: a model and seven principles of good feedback practice. *Stud. High. Educ.* **31**, 199–218 (2006).
8. Institute of Medicine (US) Committee on Clinical Practice Guidelines. *Guidelines for Clinical Practice: From Development to Use*. (National Academies Press (US), Washington (DC), 1992).
9. Rotter, T., Jong, R. B. de, Lacko, S. E., Ronellenfitsch, U. & Kinsman, L. Clinical pathways as a quality strategy. in *Improving healthcare quality in Europe: Characteristics, effectiveness and implementation of different strategies [Internet]* (European Observatory on Health Systems and Policies, 2019).
10. Terenziani, P., Montani, S., Bottrighi, A., Molino, G. & Torchio, M. Applying artificial intelligence to clinical guidelines: the GLARE approach. *Stud. Health Technol. Inform.* **139**, 273–282 (2008).
11. SNOMED CT Concept Model | Practical Guides | SNOMED International Documents. <https://docs.snomed.org/snomed-ct-practical-guides/snomed-ct-starter-guide/6-snomed-ct-concept-model> (2025).
12. Pesquita, C., Faria, D., Falcão, A. O., Lord, P. & Couto, F. M. Semantic Similarity in Biomedical Ontologies. *PLoS Comput. Biol.* **5**, e1000443 (2009).
13. Resnik, P. Semantic Similarity in a Taxonomy: An Information-Based Measure and its Application to Problems of Ambiguity in Natural Language. *J. Artif. Intell. Res.* **11**, 95–130 (1999).
14. Jiang, J. J. & Conrath, D. W. Semantic Similarity Based on Corpus Statistics and Lexical Taxonomy. (1997) doi:10.48550/ARXIV.CMP-LG/9709008.
15. Seco, N., Veale, T. & Hayes, J. An Intrinsic Information Content Metric for Semantic Similarity in WordNet. In *ECAI 2004: Proceedings of the 16th European Conference on Artificial Intelligence* (eds. López de Mántaras, R. & Saitta, L.) 1089–1090 (IOS Press, 2004).
16. Kuhn, H. W. The Hungarian method for the assignment problem. *Nav. Res. Logist. Q.* **2**, 83–97 (1955).
17. Munkres, J. Algorithms for the Assignment and Transportation Problems. *J. Soc. Ind. Appl. Math.* **5**, 32–38 (1957).

18. Transitive Closure Files | Specifications SNOMED CT Release File Specification | SNOMED International Documents. <https://docs.snomed.org/snomed-ct-specifications/snomed-ct-release-file-specification/component-release-file-specification/4.2-file-format-specifications/4.2.5-transitive-closure-files> (2025).
19. Shute, V. J. Focus on Formative Feedback. *Rev. Educ. Res.* **78**, 153–189 (2008).
20. Select Edition and Version | Practical Guides SNOMED CT Terminology Services Guide | SNOMED International Documents. <https://docs.snomed.org/snomed-ct-practical-guides/snomed-ct-terminology-services-guide/4-terminology-service-types/4.1-select-edition-and-version> (2025).