

Лекция 1. Научаване на понятия и наредба „от-общото-към-специфичното”

1.1. Увод

Една от разликите между хората и компютрите е, че обикновено хората, извършвайки някаква дейност, се стремят да подобрят начина, по който я изпълняват. С други думи, човешката дейност е неразривно свързана с процеса на самообучение, докато повечето от сегашните компютри са само изпълнители на възложени им процедури. Компютрите могат да работят много ефективно, но не и да подобряват сами своето поведение с натрупване на опита.

Изследванията в областта на *машинното самообучение* засягат създаването на компютърни програми, способни да придобиват нови знания или да подобряват вече придобити знания чрез използване на постъпващата информация. В момента тази информация (примери, факти, описания и т. н.) най-често се въвежда от човека — учител. Обаче вече съществуват и програми, използващи методи на машинно самообучение, които могат да получават данни директно от околната среда чрез разнообразни сензорни устройства или от глобалната информационна мрежа.

Бързият прогрес в машинното самообучение през последните години го превърна в един от централни фактори за развитие на всички дялове на изкуствения интелект. В частност, работите в областта на машинното самообучение са много важни за разработката на експертни системи, компютърно зрение и разпознаване на говор, автономни роботи, концептуален анализ на бази данни и интелигентни обучаващи системи и т.н. Следователно създаването на мощни самообучаващи се системи може да доведе до появата на непредсказуемо широк спектър от нови приложения.

Най-общо, една самообучаваща се система е предназначена да научава описание на определено понятие чрез използване на зададено множество от негови примери, представени от учителя, и чрез достъпни на системата основни знания за проблемната област. Примерите на понятието могат да бъдат „положителни” (например „куче” при научаване на понятието „бозайник”) или „отрицателни” (например „скорпион” за същото понятие). Основните знания за областта съдържат информация за езика, използван за описание на примерите и понятията. Те могат да включват възможните стойности на характеристиките (атрибути) на понятието и техните йерархии, предикати, допълнителни синтактични правила и т.н. Един алгоритъм за самообучение се изгражда в зависимост от типа на наличните примери, от обема и релевантността на достъпни основни знания, от начина на представяне на примери и знания, от използваните предположения за природата на понятията, които трябва да бъдат научени, и от опита на създателя на алгоритъма.

Много важно изискване към всяка самообучаваща се система е тя да може да се справя с непълни данни. Примерите често съдържат определено количество „шум“ — грешки в техните описания или класификации. Например, като класификационна грешка може да се разглежда определянето на делфин като риба от един разсеян учител. Освен това, примерите могат да бъдат непълни в смисъл, че някои атрибути стойности в тях са неизвестни (липсват). Също така основните знания за областта не винаги са пълни и точни.

По вида на създаваното описание на понятия съществуващите алгоритми за машинно самообучение могат да бъдат групирани в две основни категории — неявни (субсимволни) и явни (символни). Първите създават свои собствени представяния на научаваните понятия, които могат ефективно да се използват за цели на *разпознаването* на понятия. Обаче това *вътрешно-системно* описание на понятия не може да бъде лесно интерпретирано от потребителя (човека) и нито обяснява, нито подсказва начина, по който протича процесът на разпознаването. Субсимволните методи обикновено се основават на пресмятането на различни коефициенти, разстояния или тегла.

Една важна характеристика на символните методи за машинно самообучение е, че създаваните от тях знания (описания на понятия) по-лесно се интерпретират и се разбират от хората. Това означава, че те се опитват да направят представянето на знания по-просто от гледна точка на използваните структури. Този критерий понякога се нарича *принцип на разбираемостта*. За представяне на знания символните самообучаващи се системи обикновено използват такива структури като логически описания, правила, класификационни дървета, семантични мрежи, уравнения, фрейми и граматики.

Различието между субсимволните и символните методи за машинно самообучение може да бъде формулирано в термините на три критерия. Тези критерии — слаб, силен и свръх силен — се различават по техните изисквания към степента на разбираемост на научавани описания.

Слаб критерий. Системата използва наличните данни, за да генерира едно подлежащо на обновление описание на проблемната област, позволяващо поведението на системата да се подобрява при постъпване на нови данни.

Силен критерий. Слабият критерий е изпълнен. Освен това системата може да представя изменения на изграденото вътрешно представяне на проблемната област *в явна символна форма*.

Свръхсилен критерий. Слабият и силният критерии са изпълнени. Освен това системата може да представя изменения на вътрешното представяне *в операционно-ефективна символна форма*.

Всеки метод за машинно самообучение отговаря на изискванията на слабия критерий. На една група субсимволни методи, научаващи *екстенционалното описание* на понятия, е посветена частта на курса, оглавена *Самообучение чрез*

запомняне. В нея подробно са разгледани алгоритми, конструиращи понятия чрез избор на неговите екземпляри — примери, представящи понятието по „най-добрия начин“. Класификацията на нов обект като пример за наученото понятие се основава на измерване на неговото сходство с един или няколко екземпляра от понятието.

Основни символни методи за машинно самообучение, отговарящи на силния критерий, са разгледани в първата част на курса - *Самообучение чрез индукция*. Характерно за тези методи е, че чрез определена процедура за обобщаване на примери те създават *интенционално описание* на научаваното понятие, представляващо необходими и достатъчни условия, на които трябва да отговаря всеки обект, за да бъде класифициран като пример на това понятие. Обсъдени са основните индуктивни алгоритми, начините за оценка на индуктивните самообучаващи се системи, както и някои методологически проблеми на индукцията.

Последният, свръхсилен критерий изисква потребителят не само да разбира изведеното от метода описание на понятие, но и да може да *използва* това понятие без помощта на компютъра. С други думи, потребителят да може да изпълнява всички необходими пресмятания по прилагането на изведеното описание сам в собствената си глава.

Всички методи за машинно самообучение, споменати дотук, се основават на примери, предварително класифицирани от учител. Те се наричат методи за *направлявано* (supervised) самообучение или *самообучение с учител*. Методите за ненаправлявано (unsupervised) самообучение често наричат *самообучение без учител*. Основната задача на такива методи е създаване на таксономии от понятия от неклассифицирани обекти, като в общия случай нито самите понятия, нито техния брой са известни предварително. В основата на тези методи лежи идеята, че понятията групират *сходни* обекти, като сходството на обекти, принадлежащи към едно и също понятие, е по-голямо от сходството между обекти от различни понятия. Подробно тези методи ще бъдат разгледани в курса „Извличане на закономерности от данни“.

Машинното самообучение е едно от най-бързо развиващите се направления на изкуствения интелект и неговите методи намират все по-широко практическо приложение в такива съвременни области като техническата и медицинската диагностика, системи за подпомагане на взимането на решения, планиране, икономически и политически анализ, разпознаване на образи и говор, извличане на закономерности от данни и т. н. Последните изследвания са насочени към разработване на самообучаващи се системи, способни да комуникират с хората на много високо ниво, използвайки за тази цел мултимедийните представяния на данни и знания. Такива системи са в състояние да обясняват какво са научили и какво още трябва да научат. Това означава, че бъдещите самообучаващи се системи трябва да имат възможности за разсъждения на ниво метазнания, да генерират сами задачи за обучение, да задават въпроси, да знаят как да използват различните отговори и да осъществяват причинно-следствени и основани на модели разсъждения. С други думи, развитието на методите за машинно самообучение все повече и повече ще ни

приближава към разгадаването на удивителния феномен — способността на човека към самообучение.

1.2. Задача за научаване на понятия

Задача за формиране (научаване) на общи описания на понятия от техни конкретни (обучаващи) примери е централна за МС. В тази лекция ще разгледаме задача за научаване на понятия (concept learning) – извеждане на дефиниция на някое общо понятие (Булева функция), на база на определен набор (извадка) от положителни (т.е. маркирани като принадлежащи към това понятие) и отрицателни (не принадлежащи на понятието) примери на това понятие.

Като пример, да разгледаме задачата за научаване на целево понятие “дни, в които моят приятел Христо предпочита да спортува” (да тренира уинд-сърфинг). В Таблица 1.1 е представена една извадка от дни, в които той спортува и не спортува, описани чрез някои свои характеристики (*атрибути*). Атрибут *Харесва*, описва дали Христо харесва или не харесва да сърфира в този конкретен ден. Нашата задача е да се научим да предсказваме стойностите на този атрибут за всеки, произволен ден, използвайки за целта известните стойности на други негови атрибути.

Пример	Небе	Въздух	Влажност	Вятър	Вода	Прогноза	Харесва
1	Слънце	Топъл	Нормална	Силен	Топла	Същото	Да
2	Слънце	Топъл	Висока	Силен	Топла	Същото	Да
3	Дъжд	Студен	Висока	Силен	Топла	Промяна	Не
4	Слънце	Топъл	Висока	Силен	Студена	Промяна	Да

Табл. 1.1 Положителни и отрицателни обучаващи примери на целево понятие *Харесва*

Първо, което трябва да решим, е как да представим нашата *хипотеза* – т.е. описание на целевото понятие. Тъй като *език за описание на примерите* на търсеното целево понятие ние използвахме езика на атрибутивната логика – конюнкции от двойки атрибут – стойност, то нека в качеството на *език за представяне на хипотези* да изберем едно просто представяне, в което всяка хипотеза се представя като конюнкция от ограничения на атрибути, описващи примери на научаваното понятие. По-точно, ще представяме всяка хипотеза като вектор от шест ограничения, указващи конкретни значения на избраните от нас шест атрибута: *Небе*, *Въздух*, *Влажност*, *Вятър*, *Вода* и *Прогноза*. За всеки от тези атрибути хипотезата ще:

- указва чрез знак $?$, че всяка стойност е приемлива за този атрибут, или
- задава едно конкретно значение (например *Топла*) на атрибута, или
- указва чрез знак $\emptyset$, че нито една от възможни стойности не е приемлива

Ако някой пример x удовлетворява всички ограничения на хипотезата h , то ще приемаме, че h класифицира x като положителен пример ($h(x) = 1$).

Например, хипотезата, че Христо обича да сърфира само в студените дни с голяма влажност (не зависимо от стойности на други атрибути), ще се запише по следния начин: $\langle ?, \text{Студен}, \text{Висока}, ?, ?, ? \rangle$

При избрания начин на представяне, най-общата хипотеза е, че всеки ден е положителен пример на научаваното понятие - $\langle ?, ?, ?, ?, ?, ? \rangle$, а най-специфичната – $(\emptyset, \emptyset, \emptyset, \emptyset, \emptyset, \emptyset)$.

И така, задачата за научаване на понятие *Харесва* изисква научаване на множеството от дни, за които *Харесва* = да, описвайки това множество като конюнкция от ограничения, наложени върху атрибутите на примери. По формално:

Дадено:

- **Примери X :** възможни дни, описани със следните атрибути:
 - *Небе* (с възможни стойности *Слънце*, *Облаци* и *Дъжд*)
 - *Въздух* (със стойности *Топъл* и *Студен*)
 - *Влажност* (със стойности *Нормална* и *Висока*)
 - *Вятър* (със стойности *Силен* и *Слаб*)
 - *Вода* (със стойности *Топла* и *Студена*)
 - *Прогноза* (със стойности *Същото* и *Промяна*)
- **Хипотези H :** всяка хипотеза h се описва като конюнкция от ограничения на атрибутите *Небе*, *Въздух*, *Влажност*, *Вятър*, *Вода* и *Прогноза*. Ограниченията могат да бъдат “?” (всяка стойност допустима), “ $\emptyset$ ” (никаква стойност не е допустима) или някоя конкретна стойност.
- **Целево понятие c :** *Харесва*: $X \rightarrow [0, 1]$
- **Обучаващи примери D :** Положителни и отрицателни примери на целевото понятие (виж. Табл. 1-1)

Да се намери:

- Хипотеза h от H , такава че $h(x) = c(x)$ за всички x от X .
-

*Табл. 1-2. Задача за научаване на понятие *Харесва**

Обозначения

В общия случай, всяка задача за научаване на някое понятие може да бъде описана чрез множество от примери X , на които това целево понятие е определено, чрез целевата функция c (например, c : *Харесва*: $X \rightarrow [0, 1]$, т.е. $c(x)=1$ ако *Харесва*= да и $c(x)=0$ ако *Харесва* = не), чрез множество от хипотези H , разглеждани в процеса на обучение, и от множество от налични обучаващи примери D .

Един *обучаващ* пример се състои от примера x от X , заедно със известно значение на целевия атрибут $c(x)$. Примери, за които $c(x)=1$ се наричат *положителни*, а за които $c(x)=0$ – *отрицателни* (т.е. не принадлежащи към понятието). Обучаващите

примери ще записваме като наредена двойка $\langle x, c(x) \rangle$, а цялото им множество ще обозначаваме с D .

Знаейки множеството от обучаващи примери на целево понятие c , нашата задача е да построим хипотеза – т.е. да направим някаква приблизителна оценка за c . Чрез H ще означаваме множество от всички възможни хипотези, които можем да построим за вида на c . В общия случай, всяка една хипотеза h от H представлява някоя двоична (Булева) функция, определена на X , т.е. $h: X \rightarrow [0, 1]$. Целта е да намерим такава хипотеза, за която $h(x) = c(x)$ за всички x от X .

Основно предположение на индуктивното самообучение

Обърнете внимание, че макар задачата на обучението е да определим такава хипотеза h , която съвпада с целевото понятие c на *цялото множество от примери*, единствената информация за c е стойностите ѝ върху *множество от обучаващи примери*. Следователно алгоритмите за индуктивното самообучение в най-добрия случай могат да гарантират, че построените от тях хипотези ще съвпадат с целевото понятие върху множество от наблюдаваните обучаващи примери. Ако нямаме никаква друга информация, нашето основно предположение е, че тази хипотеза, която най-добре съвпада с обучаващите примери, ще бъде най-добра (т.е. най-близка до целевото понятие) и за *неизвестни до момента примери*.

По-формално, това *фундаментално за индуктивното самообучение предположение* може да се формулира по следния начин:

Всяка хипотеза, апроксимираща добре целевата функция върху достатъчно голямо множество от обучаващите примери, също така ще апроксимира добре тази функция и върху други, ненаблюдавани до сега примери.

1.3. Научаване на понятие като процес на търсене

Процесът на научаване на едно понятие може да се разглежда като един вид търсене в голямо пространство от хипотези, неявно зададено от представяне (езика), избрано за описание на хипотези. Целта на търсенето е да се намери хипотезата, която най-добре съвпада (приляга към) с обучаващи примери. Много важно е да отбележим, че избирайки начина за представяне на хипотези, ние неявно определяме пространството на всички възможни хипотези, които една програма за научаване може да си представи, и следователно, които *може изобщо да научи*.

Да разгледаме нашия пример с научаване на понятие *Харесва*. Имайки пред вид, че атрибутът *Небе* може да приема три възможни стойности, а останалите атрибути (*Въздух*, *Влажност*, *Вятър*, *Вода* и *Прогноза*) – по две, пространството на примери

X съдържа точно $3*2*2*2*2*2 = 96$ различни примера. Подобни изчисления показват, че H съдържа точно $5*4*4*4*4*4 = 5120$ синтактически различни хипотези. Обаче, тъй като всяка хипотеза, съдържаща един или повече символи $\emptyset$, представя празното множество от примери (т.е. класифицира всеки пример като отрицателен), пространството от семантически различни хипотези е само $1 + 4*3*3*3*3*3 = 973$. Виждаме, че нашият пример е много лесна задача за обучение с относително малко и крайно пространство на хипотези. Повече от реални задачи за научаване на понятия се описват с много голямо, а често и безкрайно пространство от хипотези.

Ако разгледаме задачата за научаване на понятия като задачата за търсене, е съвсем естествено да проучим различни стратегии за търсене в пространството на хипотезите. От практическия интерес е да проучим алгоритми, които могат да осъществяват ефективното търсене в много големи или дори безкрайни пространства, за да намират хипотезите, които най-добре съвпадат с обучаващите данни.

Наредба на хипотези от-общото-към-частното

Повечето от алгоритмите за научаване на понятия организират търсене в пространството от хипотези, базирайки се на една много полезна структура, съществуваща за всяка задача от тази област – подреждане на хипотези от общото към частното (специфичното). Използвайки тази естествена структура е възможно да се проектира алгоритъм за обучение, който е способен на изчерпващо претърсване на дори безкрайни пространства от хипотези без явното изброяване на всички хипотези. За илюстрация на това подреждане, да разгледаме две хипотези:

$$h_1 = \langle \text{Слънце}, ?, ?, \text{Силен}, ?, ? \rangle \text{ и } h_2 = \langle \text{Слънце}, ?, ?, ?, ?, ? \rangle$$

Да разгледаме сега множествата от примери, които се класифицират като положителни от хипотезите h_1 и h_2 . Тъй като h_2 налага по-малко ограничения на примери, тя класифицира повече примери като положителни. На практика, всеки пример, класифициран като положителен от h_1 , ще се класифицира като положителен и от h_2 . По тази причина, ще казваме, че хипотеза h_2 е *по-обща-от* h_1 . За да дадем по-формално определение на това отношение “*по-обща-от*”, ще казваме, че всеки пример x от X *удовлетворява* хипотеза h от H , ако $h(x) = 1$.

Ще дефинираме отношение *по-обща-или-равна-на* в термините на множества от примери, удовлетворяващи две хипотези: при зададени две хипотези h_j и h_k , ще казваме, че h_j е *по-обща-или-равна-на* h_k , тогава и само тогава, когато всеки пример, удовлетворяващ h_k също така удовлетворява и h_j .

Определение 1: Нека h_j и h_k са две Булеви функции, определени на X . Тогава h_j е *по-обща-или-равна-на* h_k ($h_j \geq_g h_k$), тогава и само тогава, когато

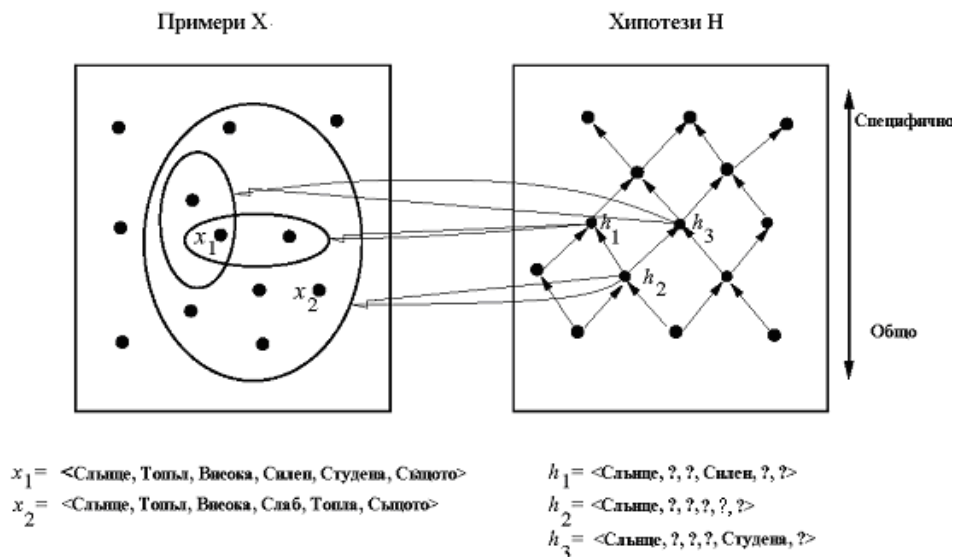
$$(\forall x \in X)[(h_k(x) = 1) \rightarrow (h_j(x) = 1)]$$

Определение 2: Нека h_j и h_k са две Булеви функции, определени на X . Тогава h_j е *по-обща-от* h_k ($h_j >_g h_k$), тогава и само тогава, когато $(h_j \geq_g h_k) \wedge (h_k \not\geq_g h_j)$

Понякога, ще използваме и обратното отношение – ще казваме, че h_j е *по-специфична-от* h_k , ако h_k е *по-обща-от* h_j .

За илюстрация на тези определения да разгледаме три хипотези h_1 , h_2 и h_3 от нашия пример с понятието *Харесва* (виж. Фиг. 1-1). Как тези хипотези се отнасят относително релацията $\geq_g$? Както се вижда, хипотезата h_2 е *по-обща-от* h_1 . Аналогично, хипотезата h_2 е *по-обща-от* h_3 . Нищо не може да се каже за отношението между h_1 и h_3 , макар, че множествата от примери, удовлетворяващи тези хипотези се пресичат.

Обърнете внимание, че релациите *по-обща-или-равна-на* и *по-обща-от* се определят независимо от целевото понятие. Те зависят само от това, кои примери ще удовлетворяват или не тези хипотези. Формално, релацията *по-обща-или-равна-на* дефинира *частичната наредба* на хипотезите в H (тази релация е рефлексивна, антисиметрична и транзитивна). Неформално, една структура е *частична наредба* (в отлечието от *пълна*), ако могат да се намерят такива двойки хипотези, например h_1 и h_3 , че $h_1 \not\geq_g h_3$ и $h_3 \not\geq_g h_1$. Релацията *по-обща-или-равна-на* е важна, тъй като дава един удобен начин за структуриране на пространството от хипотези за задача за научаване на *произволно* понятие, което може ефективно да се използва при търсене на хипотези, които най-добре пасват на обучаващите данни.



Фиг. 1-1. Левия квадрат представя множество от всички примери, а десния – от всички хипотези. На всяка хипотези съответства определено подмножество от примери – тези, които тя класифицира като положителни. Стрелките, свързващи хипотезите, представляват отношението *по-обща-от*, като стрелкичка сочи на по-малко обща хипотеза. Обърнете внимание, че подмножеството от примери, характеризиращо h_2 , включва подмножеството, характеризиращо h_1 , следователно h_2 е *по-обща-от* h_1 .

1.4. Алгоритъм FIND-S: намиране на максимално специфична хипотеза

Да разгледаме един възможен начин за използване на частичната наредба *по-обща-от*. Идеята е, да започнем с най-специфичната хипотеза в H , а след това да опитваме последователно да обобщаваме такава хипотеза всеки път, когато се окаже, че тя не покрива някакъв положителен обучаващ пример (“не покрива” означава, че хипотеза не може да класифицира този положителен обучаващ пример като такъв). Ще наречем този алгоритъм FIND-S:

Алгоритъм FIND-S:

1. Инициализирай h с възможно най-специфична хипотеза в H
 2. За всеки положителен пример x направи:
 - За всяко ограничение на атрибута a_i в h
Ако ограничението a_i е удовлетворено от x
То не прави нищо
Иначе замени a_i със следващо по-общо ограничение, което е удовлетворено от x
 3. Изведи като резултат хипотеза h
-

Да разгледаме по-подробно как работи този алгоритъм, когато последователност от обучаващи примери се задава от Таблица 1.1.

Първата стъпка е да инициализираме h с най-специфичната хипотеза в H :

$$h \leftarrow \langle \emptyset, \emptyset, \emptyset, \emptyset, \emptyset, \emptyset \rangle$$

Да разгледаме първия пример от Таблицата 1.1. – той е положителен; виждаме, че наложените от h ограничения на атрибутите на примера ($\emptyset$) са прекалено специфични – примерът не ги удовлетворява. Следователно, всяко от тези ограничения трябва да бъде заместено от следващо по-общо ограничение, което е удовлетворено от този пример. Очевидно е, че това ограничение е стойността на съответния атрибут в разглеждания обучаващ пример:

$$h \leftarrow \langle \text{Слънце}, \text{Топъл}, \text{Нормална}, \text{Силен}, \text{Топла}, \text{Същото} \rangle$$

Новата хипотеза пак е прекалено специфична – тя твърди, че всички примери са отрицателни, за изключение на един конкретен наблюдаван обучаващ пример.

При постъпване на следващия пример от Таблица 1.1. (който е пак положителен), отново възниква необходимост да построим по-обща хипотеза, която може да покрива тези два примера. Единственият начин да направим това, е да заменим със знак “?” (т.е. всички стойности са приемливи) тези ограничения на атрибути,

наложени от текуща хипотеза, които не се удовлетворяват от стойностите на атрибути в новия пример. Получената нова по-обща хипотеза придобива следния вид:

$$h \leftarrow \langle \text{Слънце}, \text{Топъл}, ?, \text{Силен}, \text{Топла}, \text{Същото} \rangle$$

Какво става при постъпване на новия пример - третия пример от таблицата? Нищо, тъй като примерът е *отрицателен*! На практика, виждаме, че алгоритъм FIND-S ще *игнорира всеки отрицателен пример*! Защо? Никакви промени не се налагат, тъй като текущата хипотеза вече правилно класифицира всеки отрицателен пример като отрицателен! В общия случай, ако предполагахме, че пространството на хипотези H действително съдържа хипотезата, описваща истинското целево понятие, и че обучаващите данни не съдържат никакви грешки, на нас никога няма да се налага да променяме текуща хипотеза h при постъпване на произволен отрицателен пример.

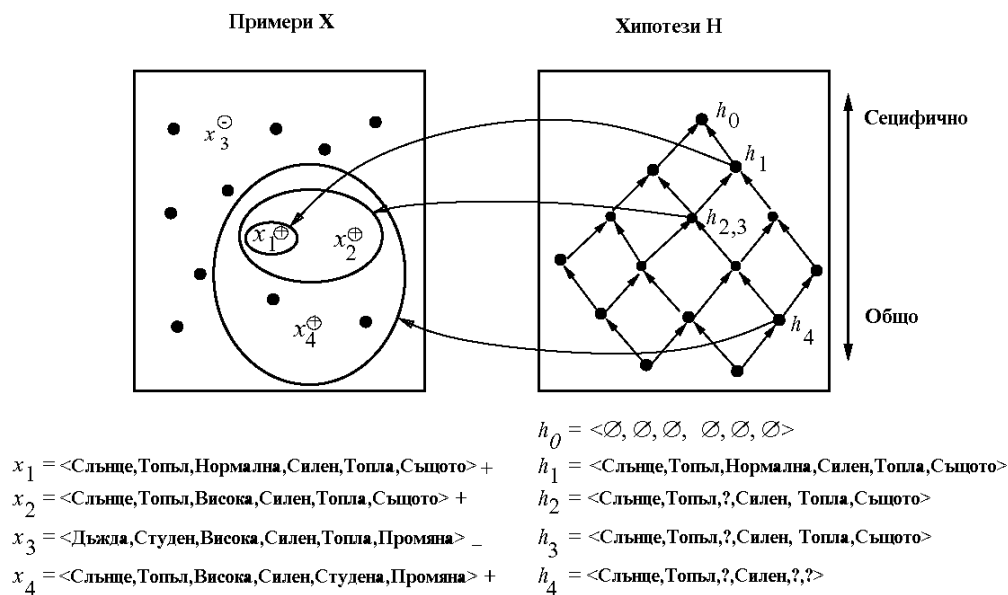
Да докажем това твърдение. Първо, да припомним, че текущата хипотеза h е най-специфичната хипотеза от H , която се намира в съгласие (не противоречи) на наблюдавани (до момента) положителни обучаващи примери на търсеното понятие. Тъй като предполагахме, че целевото понятие c също се намира в H и трябва да бъде в съгласие с положителни обучаващи примери, това означава, че c трябва да бъде *по-обща-или-равна-на* h . Обаче, тъй като целевото понятие c никога няма да покрива никакъв отрицателен пример, то същото ще важи и за h (по определение на релация *по-обща-от*). Следователно, не е необходима никаква промяна на h при постъпване на нов отрицателен пример.

Да продължим с работата на FIND-S: постъпване на четвъртия (положителен) пример от Таблица 1.1. ще доведе до следното обобщаване на текущата хипотеза:

$$h \leftarrow \langle \text{Слънце}, \text{Топъл}, ?, \text{Силен}, ?, ? \rangle$$

И така, FIND-S илюстрира един възможен начин за използване на релацията *по-обща-от* за организиране на търсене на някоя приемлива хипотеза. В процеса на търсене ние преминаваме от една хипотеза към друга по верижката от частичната наредба. На всяка стъпка текущата хипотеза се обобщава само толкова, колко да покрие нов положителен пример. По тази причина на всяка стъпка текущата хипотеза е най-специфичната хипотеза, непротиворечаща на всички разгледани до момента положителни обучаващи примери.

Основното свойство на алгоритъма FIND-S е, че в пространството на хипотези, описани чрез конюнкции от ограничения на атрибути, той гарантира намирането на най-специфичната хипотеза в H , която не противоречи на зададени положителни обучаващи примери. Неговата крайна хипотеза също така ще се съгласува с отрицателните обучаващи примери при условие, че целевото понятие се съдържа в H и обучаващите примери са коректни (т.е. не съдържат грешки).



Фиг. 1-2, Търсенето се започва с най-специфичната хипотеза (h_0), след това последователно разглежда по общи хипотези (от h_1 до h_4) в съответствие с подреждане на обучаващи примери.

Обаче, има няколко въпроса, на които алгоритъмът не дава отговори:

- *Дали алгоритъмът гарантира сходимостта към правилното целево понятие?* Макар, че FIND-S намира хипотезата, съвместима с обучаващите данни, той по никакъв начин не определя, дали намерената хипотеза е *единствено възможна* хипотеза, съвместима с тези данни (т.е. е правилното целево понятие), или съществуват и много други непротиворечиви хипотези. Бихме предпочели да имаме алгоритъм, който е способен да определи сходимостта си, или, в краен случай, да характеризира неопределеността си по отношение на истинската идентичност на целевото понятие.
- *Защо се предпочита най-специфичната хипотеза?* Когато има няколко хипотези, съвместими с обучаващите примери, FIND-S намира най-специфичната от тях. Не е ясно, дали ние трябва да предпочитаме тази хипотеза пред друга, речем, най-общата или от някакъв междинен порядък на общността.
- *Дали обучаващи примери са непротиворечиви?* В повечето от реални задачи обучаващите примери съдържат (или могат да съдържат) грешки или “шум”. Подобни противоречиви множества от обучаващи примери могат много сериозно да “излъжат” FIND-S, тъй като той не обръща внимание на отрицателните примери. За предпочитане е да имаме алгоритъм, който, най-малко, може да определя, дали обучаващите данни са противоречиви, и да може да се оправя с тях.

- *Какво да се прави, ако има няколко най-специфични хипотези? В този случай FIND-S трябва да бъде разширен с възможността да се връща за преглед на други възможности за обобщаване на хипотези и да позволява преглед на други клонове, ако се окаже че целевото описание не се намира в претърсваната част от пространството на хипотези.*