
Verificación

2.1. Introducción

En el capítulo anterior hemos visto cómo utilizar predicados para definir conjuntos de estados y hemos especificado algoritmos mediante dos predicados llamados precondition y postcondition. El primero describe el conjunto de estados válidos al comienzo del algoritmo y el segundo el conjunto de estados alcanzables a su terminación. Ahora vamos a aplicar la misma idea a la descripción del conjunto de estados alcanzados por el algoritmo en un punto intermedio de su ejecución.

Escribiendo una aserción entre cada dos instrucciones elementales del algoritmo obtenemos un modo de razonar sobre la corrección de un programa. Comenzaremos considerando un algoritmo A de la forma $A_1 ; \dots ; A_n$ (siendo cada A_i una instrucción elemental) cuya especificación es $\{P\} A \{Q\}$. Al introducir predicados intermedios, obtenemos

$$\{P \equiv R_0\} A_1 \{R_1\} ; \dots ; \{R_{n-1}\} A_n \{R_n \equiv Q\}.$$

Si el predicado inicial P se satisface y cada “algoritmo” elemental A_i consistente en una sola instrucción satisface la especificación $\{R_{i-1}\} A_i \{R_i\}$, entonces se satisface finalmente la postcondition Q y el algoritmo completo es correcto (cumple su especificación).

Necesitamos, por tanto, reglas para decidir si una instrucción elemental A satisface una especificación dada $\{P\} A \{Q\}$. En este capítulo vamos a ver dichas reglas para cada una de las instrucciones de un pseudocódigo imperativo con el que escribiremos nuestros algoritmos.

Conocer dichas reglas para cada instrucción del lenguaje es equivalente a conocer la *semántica* de este; de hecho, el conjunto de tales reglas se denomina *semántica axiomática* del lenguaje. Dichas reglas formalizan el efecto que, intuitivamente, sabemos tienen las instrucciones del lenguaje. A

28 Algoritmos correctos y eficientes

estas reglas las llamaremos *reglas de verificación* porque nos permiten *verificar* la corrección de un algoritmo con respecto a su especificación. Tendrán la siguiente forma:

$$\frac{\text{Premisas}}{\text{Conclusión}}$$

Antes de ver las reglas para cada una de las instrucciones, proporcionamos unas reglas de verificación básicas.

2.1.1. Reglas básicas de verificación

Estas reglas son generales, es decir, se pueden aplicar independientemente del algoritmo A que aparece en ellas.

Fortalecimiento de la precondition. El hecho de que una especificación sea satisfecha por un algoritmo nos aporta información sobre otras especificaciones que el mismo algoritmo puede cumplir, por ejemplo aquellas en las que se fortalece la precondition:

$$\frac{P' \Rightarrow P \quad \{P\} A \{Q\}}{\{P'\} A \{Q\}}$$

Debilitamiento de la postcondición. También podemos debilitar la postcondición y la especificación seguirá siendo correcta:

$$\frac{\{P\} A \{Q\} \quad Q \Rightarrow Q'}{\{P\} A \{Q'\}}$$

Conjunción en la postcondición. La siguiente regla es bastante intuitiva:

$$\frac{\{P\} A \{Q_1\} \quad \{P\} A \{Q_2\}}{\{P\} A \{Q_1 \wedge Q_2\}}$$

Si comenzando la ejecución en un estado que satisface P alcanzamos un estado que satisface Q_1 y también Q_2 , entonces dicho estado satisface trivialmente su conjunción.

Obsérvese que no es necesario proporcionar una regla para la conjunción en la precondition puesto que es un caso particular de la regla de fortalecimiento de la precondition.

Disyunción en la precondition. La siguiente regla también es intuitiva:

$$\frac{\{P_1\} A \{Q\} \quad \{P_2\} A \{Q\}}{\{P_1 \vee P_2\} A \{Q\}}$$

Si podemos alcanzar un estado que satisface Q cuando partimos de un estado que satisface P_1 o de uno que satisface P_2 , entonces también podremos alcanzarlo a partir de uno que satisface la disyunción.

Obsérvese que no es necesario proporcionar una regla para la disyunción en la postcondición puesto que se trata de un caso particular de la regla de debilitamiento de la postcondición.

2.1.2. Precondición más débil

Acabamos de ver que la precondición de una especificación correcta se puede fortalecer. Esto implica que para demostrar la corrección de una especificación $\{P\} A \{Q\}$ basta con determinar cuál es la precondición más débil R que cumple $\{R\} A \{Q\}$ y comprobar que P es más fuerte que R . Si denotamos a la *precondición más débil* de A con respecto a Q como $\text{pmd}(A, Q)$ podemos aplicar siempre la siguiente regla:

$$\frac{P \Rightarrow \text{pmd}(A, Q)}{\{P\} A \{Q\}}$$

A continuación proporcionamos para cada una de las instrucciones del lenguaje en que escribiremos nuestros algoritmos

- la regla de verificación para dicha instrucción y
- la forma de determinar la precondición más débil de la instrucción con respecto a una post-condición dada.

2.1.3. Reglas específicas

Instrucción nula. La instrucción **nada** no hace nada, como su nombre indica. El axioma que define su comportamiento es el siguiente:

$$\{P\} \text{ nada } \{P\}.$$

Los conjuntos de estados inicial y final son el mismo, lo que nos indica que la instrucción siempre termina y que su efecto sobre el estado del cómputo es nulo. Por tanto, la precondición más débil para la instrucción **nada** se obtiene de la siguiente manera:

$$\text{pmd}(\text{nada}, Q) \Leftrightarrow Q.$$

Asignación simple. La instrucción de asignación

$$x := e$$

donde x es una variable y e una expresión del mismo tipo que x , tiene el siguiente significado intuitivo: la expresión e es evaluada en el estado actual y su valor es asignado a la variable x , que pierde su antiguo valor. El axioma para la asignación es el siguiente:

$$\{\text{def}(e) \wedge Q_x^e\} x := e \{Q\},$$

donde $\text{def}(e)$ denota el predicado que determina que el valor de la expresión e está bien definido. Por ejemplo, si e es la expresión aritmética $x \text{ div } y$ donde x e y son variables enteras, el predicado $\text{def}(e)$ es $y \neq 0$ puesto que la división no está bien definida cuando el divisor es 0. Otras expresiones que típicamente pueden provocar errores en la evaluación de una expresión son aquellas en las que se accede a una posición de un vector, la cual debe estar dentro del rango del vector para que el acceso esté bien definido. En aquellos casos en

30 Algoritmos correctos y eficientes

los que sea evidente que el predicado $\text{def}(e)$ es equivalente a cierto omitiremos comentarios al respecto.

Si queremos que x satisfaga cierta propiedad Q después de ejecutarse la instrucción, dicha propiedad ha de ser satisfecha por e antes de ejecutarse aquella, dado que el valor de la expresión e antes de ejecutarse la instrucción va a ser copiado en x . Por tanto,

$$\text{pmd}(x := e, Q) \Leftrightarrow \text{def}(e) \wedge Q_x^e.$$

Asignación múltiple. La asignación múltiple modifica el valor de varias variables *simultáneamente*. La sintaxis general es

$$\langle x_1, \dots, x_n \rangle := \langle e_1, \dots, e_n \rangle$$

donde cada e_i es una expresión del tipo de la variable x_i . Al ejecutarse esta instrucción se evalúan *todas* las expresiones en el estado actual, sin modificar *ninguna* variable hasta que todas las expresiones estén evaluadas, momento en el que se asignan los correspondientes valores a cada variable.

El axioma para este tipo de asignación es

$$\left\{ \bigwedge_{i=1}^n \text{def}(e_i) \wedge Q_{x_1, \dots, x_n}^{e_1, \dots, e_n} \right\} \langle x_1, \dots, x_n \rangle := \langle e_1, \dots, e_n \rangle \{Q\}.$$

Y la precondition más débil

$$\text{pmd}(\langle x_1, \dots, x_n \rangle := \langle e_1, \dots, e_n \rangle, Q) \Leftrightarrow \bigwedge_{i=1}^n \text{def}(e_i) \wedge Q_{x_1, \dots, x_n}^{e_1, \dots, e_n}.$$

Asignación a vectores. La regla de la asignación no se puede aplicar a asignaciones de la forma

$$v[i] := e$$

tratando a la expresión $v[i]$ como si fuese una variable. Si permitiésemos aplicar la regla de la asignación de este modo, tendríamos:

- a) Enunciados verdaderos, pero imposibles de verificar (utilizando la regla de la asignación), como

$$\{i = j\} v[i] := 0 \{v[j] = 0\}.$$

- b) Enunciados falsos, pero que se pueden “verificar” (incorrectamente), como

$$\{v[1] = 2 \wedge v[2] = 2\} v[v[2]] := 1 \{v[v[2]] = 1\}.$$

La regla correcta de la asignación a vectores se obtiene interpretando $v[i] := e$ como una asignación que modifica al vector v , es decir, viendo la asignación como

$$v := \text{asig}(v, i, e)$$

siendo el vector $asig(v, i, e)$ un vector del mismo tipo que v que cumple:

$$asig(v, i, e)[j] = \begin{cases} e & \text{si } i = j \\ v[j] & \text{si } i \neq j \text{ y } j \text{ está en el rango de los índices de } v. \end{cases}$$

Para razonar acerca de la asignación $v[i] := e$ aplicamos a $v := asig(v, i, e)$ la regla de la asignación normal.

Composición secuencial. La composición secuencial de dos instrucciones A_1 y A_2 , denotada

$$A_1 ; A_2$$

se utiliza para encadenar (o “secuencializar”) cómputos. Como consecuencia, el estado final producido por A_1 se convierte en el estado inicial de A_2 . La siguiente regla expresa formalmente esta idea:

$$\frac{\{P\} A_1 \{R\} \quad \{R\} A_2 \{Q\}}{\{P\} A_1 ; A_2 \{Q\}}$$

Aparentemente es necesario inventar un predicado R que describa el conjunto de estados intermedios, alcanzados después de ejecutar A_1 y antes de ejecutar A_2 , y demostrar que se satisfacen las especificaciones de cada instrucción por separado. En la mayoría de los casos vamos a tomar como R la $\text{pmd}(A_2, Q)$, para solo demostrar $\{P\} A_1 \{R\}$ calculando a su vez la $\text{pmd}(A_1, R)$ y viendo que $P \Rightarrow \text{pmd}(A_1, R)$. Por tanto,

$$\text{pmd}(A_1 ; A_2, Q) \iff \text{pmd}(A_1, \text{pmd}(A_2, Q)).$$

Cuando tengamos más de dos instrucciones utilizaremos la misma idea, determinando aserciones intermedias y construyendo demostraciones para cada paso.

Composición alternativa. La sintaxis de la composición alternativa, o distinción de casos, es:

casos

$$\begin{array}{l} B_1 \rightarrow A_1 \\ \square B_2 \rightarrow A_2 \\ \vdots \\ \square B_n \rightarrow A_n \end{array}$$

fcasos

donde $A_1, \dots, A_n$ son instrucciones, posiblemente composiciones secuenciales, y $B_1, \dots, B_n$ son expresiones booleanas.

Se exige que se cumpla alguna de las guardas B_i y que estas sean excluyentes además de estar

32 Algoritmos correctos y eficientes

bien definidas:

$$\begin{array}{c}
 P \Rightarrow \bigwedge_{i=1}^n \text{def}(B_i) \\
 P \Rightarrow \bigoplus_{i=1}^n B_i \\
 \{P \wedge B_1\} A_1 \{Q\} \\
 \vdots \\
 \{P \wedge B_n\} A_n \{Q\} \\
 \hline
 \{P\} \text{ casos } B_1 \rightarrow A_1 \square \dots \square B_n \rightarrow A_n \text{ fcasos } \{Q\}
 \end{array}$$

El operador $\oplus$ denota la disyunción exclusiva.

La precondition más débil para la composición alternativa se calcula de la siguiente manera:

$$\begin{aligned}
 & \text{pmd}(\text{casos } B_1 \rightarrow A_1 \square \dots \square B_n \rightarrow A_n \text{ fcasos}, Q) \\
 \Leftrightarrow & \bigwedge_{i=1}^n \text{def}(B_i) \wedge \bigoplus_{i=1}^n B_i \wedge (B_1 \rightarrow \text{pmd}(A_1, Q)) \wedge \dots \wedge (B_n \rightarrow \text{pmd}(A_n, Q)).
 \end{aligned}$$

Una forma alternativa de verificar la distinción de casos consiste en demostrar

$$P \wedge B_i \Rightarrow \text{pmd}(A_i, Q)$$

para cada $i \in \{1..n\}$, además de comprobar que $P \Rightarrow \bigwedge_{i=1}^n \text{def}(B_i) \wedge \bigoplus_{i=1}^n B_i$.

Un caso particular de la distinción de casos se presenta cuando tenemos

casos

$$\begin{array}{l}
 B \rightarrow A_1 \\
 \square \neg B \rightarrow A_2
 \end{array}$$

fcasos

que escribimos como

si B **entonces**

$$A_1$$

si no

$$A_2$$

fsi

La regla de verificación para este caso es

$$\begin{array}{c}
 P \Rightarrow \text{def}(B) \\
 \{P \wedge B\} A_1 \{Q\} \\
 \{P \wedge \neg B\} A_2 \{Q\} \\
 \hline
 \{P\} \text{ si } B \text{ entonces } A_1 \text{ si no } A_2 \text{ fsi } \{Q\}
 \end{array}$$

y la precondition más débil

$$\begin{aligned} & \text{pmd}(\text{si } B \text{ entonces } A_1 \text{ si no } A_2 \text{ fsi}, Q) \\ \Leftrightarrow & \text{def}(B) \wedge ((B \wedge \text{pmd}(A_1, Q)) \vee (\neg B \wedge \text{pmd}(A_2, Q))). \end{aligned}$$

Obsérvese que este último predicado es equivalente a $(B \rightarrow \text{pmd}(A_1, Q)) \wedge (\neg B \rightarrow \text{pmd}(A_2, Q))$, obtenido calculando la precondition más débil de la distinción de casos equivalente.

Composición iterativa. La instrucción de iteración de nuestro lenguaje es la instrucción **mientras** cuya sintaxis es:

mientras B **hacer**
 A
fmientras

donde B es una expresión booleana y A una instrucción. La instrucción comienza evaluando la *condición de iteración* B : si es falsa, se finaliza la ejecución del bucle; en otro caso, se ejecuta el *cuerpo* A de la iteración y se vuelve a evaluar la condición B , repitiendo el proceso.

Para la verificación de un bucle necesitamos un predicado llamado *invariante* que describe los distintos estados por los que pasa el bucle, estableciendo la relación existente entre las variables que intervienen en este. El invariante lo escribiremos como anotación junto con el bucle:

$\{P\}$
 $\{I\}$
mientras B **hacer**
 A
fmientras
 $\{Q\}$

El invariante I tiene que cumplir las siguientes condiciones:

(i.1) Se satisface antes de empezar el bucle, es decir, antes de la primera iteración. Esto lo expresamos requiriendo que

$$P \Rightarrow I.$$

(i.2) Se mantiene al ejecutar el cuerpo A del bucle. El cuerpo se ejecuta solamente si la condición se cumple, es decir, tendrá que cumplirse

$$\{I \wedge B\} A \{I\}.$$

(i.3) Se cumple al salir del bucle, cuando B se hace falsa, y esto nos tiene que llevar a la postcondición:

$$I \wedge \neg B \Rightarrow Q.$$

34 Algoritmos correctos y eficientes

Además de estos tres puntos, tendremos que probar que la instrucción iterativa *termina*. Para ello tendremos que buscar una función C que dependa de las variables del bucle (generalmente solo de aquellas variables que aparecen en la condición de iteración B) y que tome valores enteros, de forma que se cumpla que:

(c.1) Es mayor o igual que cero cuando se cumple la condición B :

$$I \wedge B \Rightarrow C \geq 0.$$

(c.2) Decrece al ejecutar el cuerpo A del bucle:

$$\{I \wedge B \wedge C = T\} A \{C < T\}.$$

Esta función es una cota superior del número de iteraciones que quedan por realizar y por eso nos referiremos a ella como *función de cota*.

De todo lo anterior concluimos la regla de verificación para la instrucción iterativa:

$$\frac{\begin{array}{c} P \Rightarrow I \\ \{I \wedge B\} A \{I\} \\ I \wedge \neg B \Rightarrow Q \\ I \wedge B \Rightarrow C \geq 0 \\ \{I \wedge B \wedge C = T\} A \{C < T\} \end{array}}{\{P\} \text{ mientras } B \text{ hacer } A \text{ fmientras } \{Q\}}$$

2.1.4. Verificación de algoritmos recursivos

La recursividad es una característica de la mayoría de los lenguajes de programación que permite que un procedimiento o función haga referencia a sí mismo dentro de su definición. La recursividad se utiliza, al igual que la iteración, para describir cálculos que se repiten cierto número de veces.

Razonar recursivamente permite resolver un problema P sobre unos datos D suponiendo que P ya está resuelto para otros datos D' , del mismo tipo que D y en algún sentido más sencillos (caso recursivo). Para que el razonamiento pueda aplicarse se requiere que P pueda ser resuelto directamente para unos datos suficientemente simples (caso base).

Un algoritmo es *recursivo lineal* si su ejecución genera una única llamada recursiva. Un algoritmo es *recursivo múltiple* si una misma invocación genera más de una llamada recursiva.

El esquema general de un algoritmo recursivo lineal es

```
{ P( $\bar{x}$ ) }
fun f-rec( $\bar{x} : T_1$ ) dev  $\langle \bar{y} : T_2 \rangle$ 
  casos
     $B_t(\bar{x}) \rightarrow \bar{y} := \text{triv}(\bar{x})$ 
     $\square B_{nt}(\bar{x}) \rightarrow \bar{y} := c(\text{f-rec}(s(\bar{x})), \bar{x})$ 
  fcasos
ffun
{ Q( $\bar{x}, \bar{y}$ ) }
```

donde: