

지식을 이용한 영상 융합의 제어

。 정 진영*, 박 충식**, 장 문익*, 김 재희*

* : 연세대학교 기계·전자 공학부

** : 영동공과대학 전자계산학과

E-mail : jjjung@seraph.yonsei.ac.kr, jhkim@bubble.yonsei.ac.kr

Tel : (02)361-2869, Fax : (02)362-0413

Knowledge-Based Control of Image Fusion

。 Jin Young Jung*, Choong-Shik Park**, Munik Chang*, Jaihie Kiim*

* School of Mechanical and Electrical Engineering, Yonsei Univ.

** Dept. of Computer Engineering, Youngdong Institute of Technology

요 약

본 논문에서는 지식을 이용하여 주어진 상황에서 얻어진 정보 및 정보들을 융합하는 레벨을 선택하여 적용할 수 있는 지식 레벨의 제어구조를 제안하였다.

제안된 제어 구조는 태스크 개념을 도입하여 융합을 위한 문제 분야의 지식과 제어를 위한 지식을 분리하여 생각하기 쉬우며, 관심의 대상이 되는 부분에 대한 정보들만을 융합함으로써 능동적이고 효과적인 융합을 수행할 수 있다.

1. 서 론

영상 정보는 정보 융합에서 사용되는 여러 가지 정보중의 하나로 정보의 형태가 인간의 시각과 동일하거나 거의 유사하다는 점에서 다른 어떤 형태의 정보보다 그 중요성이나 가치가 더 높다고 할 수 있다[1, 2]. 영상센서는 각각 물리적 특성이 달라서 감지할 수 있는 영역에 한계가 있기 때문에 필요한 정보를 얻지 못하는 경우가 많으며, 또한 처리과정에서 발생하는 많은 오차 때문에 본질적으로 불확실한 경우가 많다. 이와 같은 단일 영상 센서의 단점을 극복하기 위해서 최근 들어서는 여러 가지 영상 센서를 복합적으로 사용하여, 단일 센서로는 얻을 수 없는 정보들을 얻기 위한 방법들이 연구되고 있다.

본 논문에서는 앞에서 언급한 것과 같은 여러 가지 레벨에서 발생하는 융합들을 선택하여 수행

할 수 있고, 경우에 따라서 정보를 선택적으로 융합할 수 있도록 하는 지식레벨 제어구조에 대하여 제안하였다. 제안된 제어 구조는 태스크 개념을 도입하여 융합을 위한 문제 분야의 지식과 제어를 위한 지식을 분리해 내기 쉬우며, 능동적이고 효과적인 융합을 수행할 수 있다.

2. 지식을 이용한 영상 융합의 개요

영상 융합은 다양한 영상 센서로부터 얻어진 영상 정보들을 융합하여, 각각의 센서로부터는 얻을 수 없는 정보를 얻어내거나, 센서들 상호간의 상승작용에 의하여 정확한 결과를 얻어내는데 그 목적이 있다. 특히, 실세계에 존재하는 장면인 경우에는 각 센서의 다양한 특성 때문에 입력 영상내의 물체에 대하여 각각 강조되는 정보가 다르다. 따라서, 영상 융합은 특히 실세계의 영상에 대한 분석이나 해석을 하는 경우에 유용하게 사용될 수 있다.

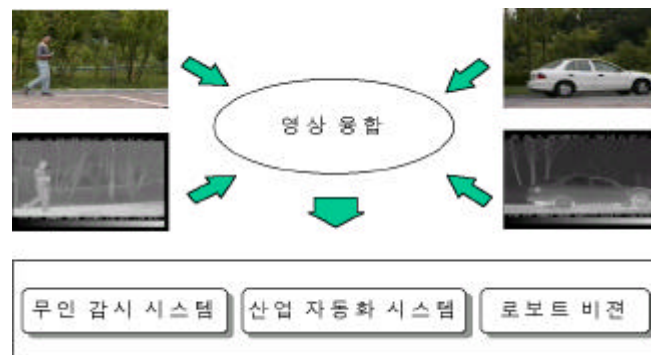


그림 1 영상 융합의 개요

현재까지 연구된 영상 융합은 그 처리단계에 따라 크게 신호 레벨, 화소 레벨, 특징 레벨, 심볼 레벨의 4단계로 구분할 수 있다[3, 4, 5]. 그러나, 융합 대상이 되는 센서 정보들이 방대하고, 다양한 형태로 나타나게 된다. 따라서, 이와 같이 방대한 정보들 중에서 어떤 정보들을 먼저 융합하고, 어떤 정보들을 이후에 융합할 것인지를 결정하는 문제는 매우 어렵고 복잡하다. 따라서, 이와 같이 방대하고, 다양한 정보를 처리하기 위해서는 지식의 사용이 필수적이다[5, 6].

3. 지식 레벨 제어에 의한 영상 융합의 구조

본 논문에서는 영상 융합을 위한 지식 기반의 제어 구조에 대하여 제안하였다. 지식 레벨에서의 제어를 위하여 태스크 개념을 도입하여 각 태스크 별로 작업 단위를 구별하였으며, 각각의 태스크의 구동에 의하여 전체적인 작업이 이루어진다.

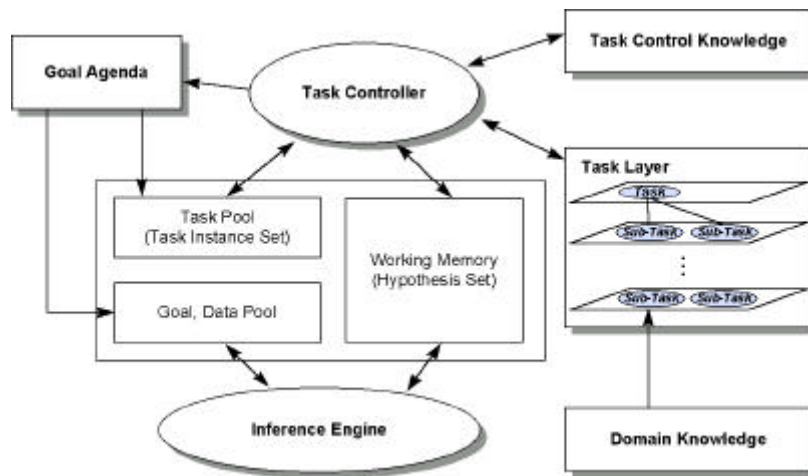


그림 2 전체 구조도

이와 같은 태스크의 수행에 관련된 행동 지침은 태스크 내의 규칙으로 표현되어지며, 각 태스크들은 상황에 따라 적절히 선정되어 결과를 내게 된다. 그림 2는 지식 레벨에서의 제어를 위한 전체 시스템의 구조도이다. 각각의 태스크는 자신이 수행하기 위한 문제 영역에 대한 지식을 규칙으로 저장하고 있으며, 태스크 컨트롤러에 의하여 구동된다. 태스크 컨트롤러가 태스크 계층 중에서 하나의 태스크를 선정하여 수행하면, 이 태스크는 태스크 풀에 들어가게 되며, 이 태스크에 의하여 얻어진 결과는 작업 메모리 공간에 놓이게 된다. 태스크의 수행은 자기 자신의 규칙과 작업 메모리 공간에 놓인 정보를 이용하여 추론 엔진을 통해 수행되며 그 결과를 다시 작업 메모리 공간에 첨가한다. 전체의 과정은 시스템의 목적이 완료될 때까지 계속해서 수행되며, 더 이상 작업 메모리 공간의 변화가 없어서 태스크의 작업이 수행될 수 없으면, 시스템의 목적은 실패로 끝난다.

3-1. 태스크의 정의

지식을 사용하기 위한 지식 기반 시스템을 구성하기 위해서는 지식을 표현하기 위한 적절한 수준의 표현단위가 정의되어야 한다. 이러한 표현단위는 인간의 사고 단위와 일치되도록 설정되고 전체 문제풀이 지식 체계 내에서 독립적으로 처리할 수 있는 단위가 되어야 한다. 그러므로 이러한 표현단위는 독자적인 목표와 이 목표를 위한 지식 및 처리기능을 가져야 한다. 또한, 이 표현

단위는 전체 지식 체계 내에서 이 표현 단위의 수행 시기, 수행에 필요한 자료의 구비 조건, 수행 결과에 따른 다른 표현단위에의 영향 등에 대한 묘사가 이루어져야 하며, 이와 같은 사고의 표현 단위를 태스크라고 정의한다[7,8].

3-2. 영상 융합을 위한 작업 단위의 분류

본 논문에서는 영상 융합을 수행하기 위한 기본적인 작업 단위들을 태스크로 분류하여 설정하였다. 이와 같은 기본적인 작업 단위들은 앞에서 정의된 정보원과 비슷하게 분류되며, 각각의 작업 단위인 태스크들은 자기 자신의 작업을 수행하기 위한 규칙과 자신이 수행하기 위하여 기본적으로 만족되어야 할 조건, 상하위 태스크들간의 종속관계 등을 포함해야 한다. 이와 같은 작업 개념 단위에서의 작업 구분은 영상 융합시 발생 가능한 모든 레벨의 융합을 수용할 수 있으며, 상호간에 연동적으로 구동됨으로써 모든 작업의 수행 과정을 파악할 수 있다.

3-2- 1. 태스크의 구성

각각의 태스크는 일의 처리를 위한 기본 작업 단위로 자신이 수행하기 위한 규칙, 자신의 부 태스크들에 대한 정보, 자기 자신의 목적, 자신의 하위 태스크를 구성하는 방법, 자신이 수행하기 위해서 먼저 수행되어야 하는 태스크 및 이 태스크의 상태 등에 대한 정보를 기술하는 슬롯을 가지고 있다. 다음은 태스크의 클래스에 대한 예이다.

```
(TaskName is-a task)
(TaskName goal (? (generate-hypothesis is-my goal)))
(TaskName state sleep)
(TaskName type rule)
(TaskName sub_ task nil)
(TaskName decomposition_ type nil)
(TaskName rule_ or_ ft (rules))
(TaskName precede generate-context)
(TaskName precondition (equal (get-list (get-list CTASK 'precede) 'state) 'success))
(TaskName repeat_ condition nil)
```

그림 3 태스크의 구성에 대한 예

여기서 이 태스크의 이름이 TaskName이고 맨 윗줄은 이러한 이름을 가진 것이 태스크의 한 종류임을 나타내며, 'goal'은 이 태스크를 수행하는 목적이다. 만약 이 태스크가 수행되어 목적을 달성하게 되면 상태가 '성공'이고 그렇지 못한 경우에는 '실패'가 된다. 모든 클래스의 상태는 처음에 수면상태이고 수행 중에 다음의 그림과 같이 하나의 태스크의 상태는 여러 가지로 변하게 된다. 종류(type)는 이 태스크를 수행하기 위한 방법에 대한 것으로 거의 대부분의 경우 규칙으로

표현되는 지식에 의하여 수행되지만 함수에 의하여 수행될 수도 있다. 만약 이 태스크의 수행방식이 규칙에 의한 것이라면 'rule_ or_ ft' 슬롯에 실제로 사용되기 위한 규칙의 이름들을 나열한다. 하위 태스크에 대한 표현은 'sub_ task' 슬롯에 기술되며 부 태스크가 없는 경우에는 'nil'로 부 태스크가 있는 경우에는 부 태스크들의 이름을 적어주며, 이러한 부 태스크들의 분할 형태는 'decomposition_ type' 슬롯에 적어준다. 부 태스크들의 분할 형태는 부 태스크들중의 하나가 목적을 달성하면 자기 자신이 수행하는 경우와 부 태스크들 모두가 목적을 달성해야 자기 자신이 수행하는 경우, 부 태스크들의 목적 달성에 상관없이 부 태스크들이 수행되지만 하면 자신이 수행하는 세 가지 종류가 있다. 선행자(precede)는 이 태스크 이전에 수행되어야 할 태스크에 대한 이름을 의미하고, 'precondition'은 이 태스크가 수행되기 위하여 만족되어야 하는 조건을 기술하는 슬롯이다. 예에서는 선행태스크의 상태가 성공이어야만 이 태스크를 수행할 수 있다. 특수한 경우로 하나의 태스크는 반복 수행 조건을 가질 수 있는데, 이 경우에는 반복 수행 조건이 만족되는 동안 계속해서 태스크의 수행이 이루어진다.

3-2-2. 영상 융합을 위한 태스크의 분류

본 논문에서 영상 융합을 적용하기 위하여 분류한 태스크의 계층도가 그림 4에 나타나 있다.

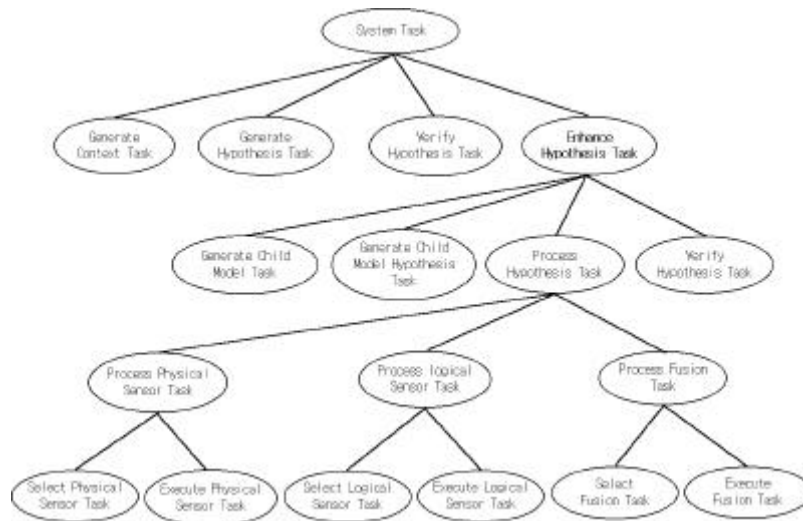


그림 4 영상 융합을 위한 태스크의 계층도

태스크는 모두 18개로 구분되어 있으며 모든 태스크는 기본적으로 앞에서 설명한 클래스와 같이 정의되어 있다. 그림에서 'enhance- hypothesis' 태스크는 반복 수행을 하는 태스크로서 4개의 부

태스크들이 모두 성공이 되거나, 모든 가능한 경우를 다 시도해보아도 부 태스크들이 성공이 되지 않는 경우에 더 이상 반복 수행을 하지 않는다.

① 시스템 태스크

시스템 태스크는 전체 시스템에 대한 입력에 대한 태스크로 사용자로부터 검출하고자 하는 대상 물체에 대한 질문을 입력받아서 전체의 태스크들을 수행하는 제어 역할을 한다.

② 환경 생성 태스크

시스템 태스크에 의하여 전체 시스템이 초기화가 되면 현재 장면에 대한 묘사를 수행한다. 현재 장면에 대한 묘사는 사용 가능한 센서들과, 센서들의 위치, 현재 시간, 묘사되는 장면이 무엇에 대한 것인지등에 대한 것들이 포함된다.

③ 가설 생성 태스크

환경 생성 태스크에 의하여 현재 장면에 대한 묘사가 이루어지면 가설 생성 태스크는 현재의 장면에 대한 지식을 이용하여 이 장면에 대해 포함되리라고 판단되는 검출 대상 물체들에 가설들을 생성한다.

④ 가설 검증 태스크

가설 생성 태스크에 의하여 가설들이 생성되면, 가설 검증 태스크는 생성된 각각의 가설들을 검증한다. 만약 정보에 의하여 가설이 검증되면, 시스템 태스크에 의하여 이 가설이 문제의 답인지 판별하게 되며, 가설이 검증되지 않으면 가설 진행 태스크로 제어를 넘겨주게 된다.

⑤ 가설 진행 태스크

가설 진행 태스크는 생성된 각각의 가설들 중에서 아직 검증이 되지 않은 가설들이 존재하는 경우, 반복 수행을 통해서 이 가설에 대한 가능성을 계산하는 태스크이다. 따라서, 검증되어야 할 가설이 존재하는 경우에는 항상 수행되며, 이 태스크는 마지막 최하위 레벨의 태스크가 수행 될 때까지 반복 수행한다. 가설 진행 태스크는 하위 모델 생성 태스크, 하위 모델 가설 생성 태스크, 가설 처리 태스크, 가설 검증 태스크와 같은 부 태스크들을 가진다.

㉠ 하위 모델 생성 태스크

하위 모델 생성 태스크는 실제로 검증해야 할 가설이 아직 검증되지 않은 경우에 이 가설에 의하여 검증되어야 하는 모델에 대한 데이터 계층상의 하위 모델의 구성 요소들에 대한 가설을 생성하는 역할을 한다.

㉠ 하위 모델 가설 생성 태스크

하위 모델 가설 생성 태스크에서 얻어진 각각의 모델 구성 요소에 대하여 가설을 생성하는 역할을 수행한다.

㉡ 하위 모델 가설 처리 태스크

실제로 생성된 각각의 모델 구성 요소에 대한 가설을 처리하기 위한 태스크로 각각의 가설에 대하여 적용할 수 있는 물리적인 센서에 대한 처리를 하기 위한 태스크와 논리적인 센서에 대한 처리를 하기 위한 태스크, 결과로 얻어진 대상에 대한 융합 처리를 하기 위한 세 종류의 부 태스크를 갖는다.

㉠ 물리적 센서 처리 태스크

물리적 센서 처리 태스크는 실제적으로 생성된 가설에서 고려하는 모델을 검증하기 위하여 환경에서 사용할 수 있는 물리적인 센서를 선택하여 수행하는 역할을 수행하며, 센서를 선택하고, 수행하는 두 가지의 부 태스크를 가지고 있다.

㉢ 하위 모델 검증 태스크

하위 모델 검증 태스크에서는 생성된 각각의 하위 모델에 대한 가설들을 검증한다. 초기상태에서는 가설들에 대한 가능성이 없으므로 검증대상으로 고려하지 않으며, 하위 모델 가설 처리 태스크에 의하여 처리된 가설들에 대해서 검증을 수행한다.

3-2-3. 태스크 계층의 종속 관계

본 논문에서 영상 융합의 처리를 위한 고안한 태스크의 계층은 최종단의 노드에 위치한 태스크를 제외하고는 모두 하위의 부 태스크를 가지고 있다. 부 태스크는 상위에 위치한 주 태스크의 작업에 필수적이며, 부 태스크가 먼저 수행되어야한 상위의 주 태스크가 작업을 수행할 수 있다. 주 태스크와 부 태스크와의 관계는 다음과 같이 세 가지의 종류로 구분된다. 첫째는 부 태스크들이 모두 성공 상태가 되어야만 주 태스크가 수행을 할 수 있는 것이다. 이것은 상위의 태스크가

하위의 태스크에서 이루어지는 작업에 매우 의존적인 것을 나타내며, 반드시 모든 부 태스크들이 성공이 되어야만 하며, 'AND' 노드에 해당한다. 둘째로 부 태스크들중에서 하나만이라도 성공이 되면 주 태스크로 제어가 넘겨지는 경우이다. 이와 같은 노드는 주 태스크를 수행하는 작업이 부 태스크에서 발생할 수 있는 정보들 중에서 하나만 있어도 수행 가능한 경우에 설정되는데, 'OR' 노드로 표현된다. 마지막하나는 부 태스크들이 자신의 목적을 달성했는지의 여부에 관계없이 단순히 수행되기만 하면 주 태스크로 제어가 넘어가는 경우이다. 이것은 'DONE'관계의 노드가 된다.

시스템 태스크는 하위의 각각의 태스크인 환경 생성 태스크, 가설 생성 태스크, 가설 검증 태스크, 가설 진행 태스크가 수행되기만 하면 다시 자기 자신이 제어를 넘겨받아 일을 수행하게 된다. 가설 진행 태스크는 현재 처리 해야하는 가설에 대해서 발생 가능한 모든 하위의 가설들에 대한 처리를 수행해야 하므로 모든 하위 가설들이 검증 될 때까지 반복 수행되며, 각각의 부 태스크를 차례로 적용하여 하위의 가설들을 차례로 검증한다. 각 계층에서의 태스크간의 선행 관계는 다음의 그림과 같다.

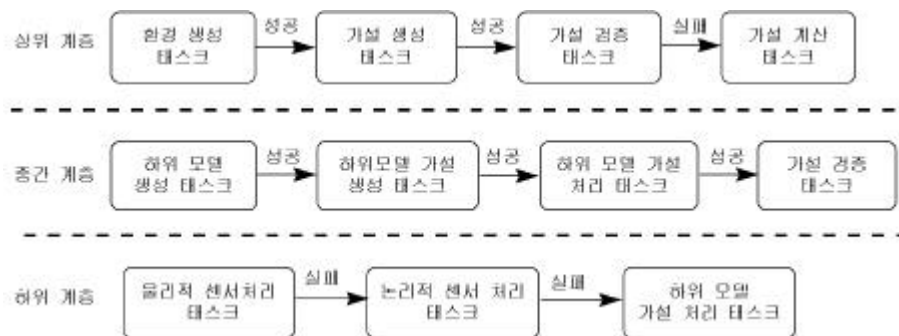


그림 5 각 계층에서의 태스크간 선행 관계

3-3. 지식 레벨 제어를 위한 지식 표현 및 선택적 융합

영상을 이용하여 영상을 해석하는 시스템은 관측되는 장면에 존재하는 대상 물체와 같은 문제 분야에 대한 선행적인 지식을 필요로 한다. 더욱이 표현되는 정보가 상이한 다양한 센서에서 입력되는 정보들을 융합하고 각 계층에서 발생하는 데이터를 적절히 사용하기 위해서는 사용 가능한 센서들의 특징 등에 관련된 사전 지식이 필수적이며 사전 지식들은 문제를 해결하기 위하여 상황에 따른 최적의 수행 방법에 대한 선택을 가능하게 한다. 모든 문제 분야에 관련된 지식은 규칙으로 표현되었으며, 각각의 태스크는 자기 자신이 수행하기 위한 별도의 지식을 포함하고 있다.

실세계의 영상들을 융합하고 영상에 존재하는 물체의 검출을 수행하기 위해서는 실제 존재하는 세계에 대한 지식이 필수적이며, 실세계에 대한 지식은 환경에 대한 지식과 대상 물체에 대한 지식, 센서에 대한 지식 등을 포함한다.

환경은 센서에 의하여 영상이 획득되기 이전의 모든 가능한 상태에 대한 묘사를 의미하는 것으로 응용 분야에 따라 각각 다른 환경에서 작동하므로 다른 환경에 대한 지식을 필요로 한다. 환경에 대한 지식은 전체 문제 분야를 해결하기 위하여 필수적이며, 지식이 없는 경우에는 아무런 사전 정보가 없는 상황에서 문제를 해결해야 하므로 가능한 모든 경우에 대해서 모든 처리를 다 해야 하는 번거로움이 있다.

대상 물체에 대한 지식은 영상내에 포함되어 있는 사람과 차량에 대한 지식을 포함한다. 각각의 물체에 대한 지식은 이 물체들이 현재 고려되고 있는 환경 하에서 센서에 의하여 감지될 수 있는 특성과 각 물체들이 갖는 구조적인 형태에 관련된 지식, 공간적인 위치에 관련된 지식등이 있다. 대상 물체의 구조적인 정보에 대한 지식은 대상 물체가 어떻게 구성되어 있으며, 물체를 구성하는 각각의 구성 요소에 대한 묘사로 이루어져 있다.

영상 융합을 위해서는 앞에서 언급한 각각의 물체가 주어진 환경에서 어떤 센서에 의하여 더 잘 감지될 수 있으며, 어떤 경우에 어떤 센서를 적용할 것인가가 정의되어 있어야 한다. 이와 같은 융합 관련 지식들은 물체나 센서, 환경에 의한 지식들을 적절히 사용하기 위한 지식이라고 생각할 수 있으며, 이러한 의미에서 지식의 처리를 위한 메타 레벨의 지식이라고 생각할 수 있다.

4. 지식을 이용한 태스크의 수행 예

지식을 이용하여 영상 융합의 과정을 처리하는 과정에 대한 예는 다음과 같다. 먼저 환경에 대한 묘사를 하는 태스크에 의하여 다음과 같은 장면에 대하여 환경에 대한 기술이 이루어진다. 현재의 장면에 대한 묘사가 이루어지면 사용자로부터 검출할 대상을 입력받아 이에 대한 가설을 만든다. 가설은 여러 종류의 속성에 의하여 표현된다. 현재의 예에서는 이 가설이 표현하는 모델이 버스로 설정되어 있고, 초기에 생성된 가설이므로 상위나 하위의 가설이 존재하지 않는다. 또한 아직 아무런 센서에 대해서도 처리되지 않은 상태이다. 가설이 생성되면 이 가설에 대한 검증을 시도하는데 현재의 경우 처음에 생성된 가설이므로 검증에 실패하게 되고, 가설 진행 태스크에 의하여 지식을 이용하여 하위의 모델에 대한 가설을 다시 생성하게 된다.

하위의 모델에 대한 가설은 위의 가설과 동일한 속성을 가지며, 다만 가설의 내용이 모델 계층 상에서 버스보다 아래에 존재하는 구성 요소이고, 상위의 가설이 버스에 대한 가설로 설정된다는 점만 다르다. 현재는 버스를 구성하는 요소중에서 몸체에 대한 가설이 생성되었다고 가정한다.

하위의 모델에 대한 가설이 설정되면 이 모델에 대한 센서의 처리를 한다. 그러나, 몸체라는 요소를 바로 추출할 수 있는 센서가 없으므로 하위 모델 가설을 처리하기 위한 태스크들을 모두 실패로 끝나게 되고, 하위 가설 검증 태스크는 선행 태스크가 실패가 되었으므로 수행되지 않고 다시 가설 진행 태스크로 제어가 넘어가게 된다. 동일한 과정에 의하여 다시 하위 모델에 대한 가설이 생성된다. 만약, 다각형에 대한 가설이 생성되었다고 가정하면, 이 가설에 대하여 센서를 적용하기 위하여 시도한다. 다각형은 논리적 센서로부터 얻을 수 있으나, 아직 영상이 획득되지 못한 상태이므로 논리적 센서를 적용하는데 실패한다. 따라서, 다시 가설 진행 태스크에게 제어가 넘어가게 되고 이 과정은 가장 하위 모델인 영상에 대한 가설이 생성될 때까지 반복된다. 영상에 대한 가설이 생성되면 이때 영상을 얻기 위하여 물리적인 센서를 적용한다. 이때 현재 가설의 대상이 영상이지만 영상의 상위 가설에는 버스가 있으므로, 어떤 물리적인 센서를 적용할 것인가는 최상위 모델인 버스에 의하여 결정된다. 먼저 물체의 검출에는 열 영상 센서가 적합하므로 열 영상 센서가 선택되어 영상이 얻어진다.

물리적 센서를 선정하기 위한 규칙 :

```
(sps-r1 if (가설이 존재하고)
    (가설의 내용이 모델 계층상의 최하위 모델이고)
    (가설의 제일 상위 모델이 버스이고)
    (이 모델이 물리적 센서에서 얻어질 수 있고)
    (버스가 물리적 센서에서 얻어질 수 있고)
    (버스에 대해 먼저 적용할 센서가 아직 적용되지 않았으면))
(sps-r1 then (먼저 적용할 센서를 선택한다)
    (다음에 적용할 센서에 대한 목록을 작성한다)
    (현재 적용한 센서에서 찾고자하는 대상물체에 대하여 기록한다))
(sps-r1 priority 90)
```

물리적 센서를 수행하기 위한 규칙

```
(eps-r1 if (물리적 센서에 대한 선정이 이루어지면))
(eps-r1 then (물리적 센서를 이용하여 영상을 획득한다) )
(eps-r1 priority 90)
```

그림 6. 물리적 센서를 구동하기 위한 규칙의 예

물리적 센서에 의하여 영상이 얻어지면 현재의 가설에 대한 검증이 이루어지게 되고 하위의 가설에 대한 검증이 이루어지면, 현재의 가설의 상위 가설에 대하여 센서를 적용하게 된다. 이와 같이 상위의 가설이 만족되거나 수행되기 위해서는 하위의 가설들이 차례대로 생성되고, 이 가설들이 만족되어야 하며, 하위의 가설이 검증되면 상위의 가설들에 대하여 차례대로 센서를 적용하고 검증해 나가는 과정을 반복해서 수행하게 된다. 동일한 과정이 반복되어 모든 가설이 검증되면 하위 모델에 대한 작업을 종료하게 된다.

지식에 의하여 융합은 모델 계층상의 한 단계에 있는 모델에 대하여 서로 다른 센서로 적용된 경우가 둘 이상 있는 경우에 수행하게 되며, 심볼 레벨을 제외하고는 결과는 영상으로 표현한다. 융합을 위한 방법으로는 2장에서 설명한 Chu의 융합 방법을 선택하여 사용하였으며, 융합을 하기 위한 기본 조건은 다음과 같다.

- 동일한 모델에 관련된 서로 다른 영상에서 얻어진 경우가 둘 이상 있는 경우
- 특정 레벨의 융합인 경우에, 지식에 의하여 각각의 모델을 구성하는 특징들을 추출하므로, 처리되는 영상들에 대한 우선 순위가 결정된다. 따라서 한번에 두 가지의 경우에 대해서만 융합한다.
- 현재 융합을 적용하기 위한 모델에 따라서 전체 영상이 될 수도, 부분 영상에 대한 정보들일 수도 있다.
- 각 경우에 대한 지식을 이용하여 융합 방법을 설정한다.

그림 7. 융합을 하기 위한 기본 조건

5. 결 론

본 논문에서는 다양한 영상 센서로부터 얻어지는 영상 정보들을 융합하는 경우에 환경과 물체 센서에 대한 지식을 사용함으로써 적절한 센서를 선택할 수 있고, 필요한 경우에 센서를 적용할 수 있는 지식 레벨에서 융합을 제어하기 위한 구조에 대하여 제안하였다. 컴퓨터 비전의 일반적인 구조에서처럼 이 구조는 서로 다른 계층을 처리할 수 있는 작업단위로 구분되어 있으며, 각각의 작업단위인 태스크의 상호관계에 의하여 전체적인 수행이 이루어진다.

본 논문에서는 지식 레벨에서 융합을 제어 할 수 있는 구조를 제안하였으며, 이 구조에서는 정보의 발생과 발생된 정보들에 대한 제어를 파악하기 쉽다. 실제로 발생할 수 있는 다양한 정보와 사용할 수 있는 정보원에 대하여 분류하였으며, 태스크 개념을 이용하여 융합을 작업단위 별로 구분하였다. 제안한 구조는 영상 융합을 하는 경우에 발생할 수 있는 모든 레벨이 하나의 반복 수행 구조 안에서 자동으로 선택될 수 있으며, 대상 물체에 대한 지식을 이용함으로써, 전체 영상에 대하여 융합을 수행하지 않고, 대상 물체에 관련된 영상의 부분만을 선택하여 융합할 수 있다. 따라서, 융합된 결과 영상을 이용하여 물체를 인식하거나 영상을 이해하는 경우에도 별도의 과정이 필요없이 결과 영상을 그대로 사용할 수 있다.

6. 참고 문헌

- [1] Richard T. Antony, *Principles of Data Fusion Automation*, Artech House, 1995.
- [2] Luo, R. C., and M. G. Kay, "Multisensor Integration and Fusion in Intelligent System",

IEEE Trans. on System Man and Cybernetics, Vol. 19, No. 5, pp.901-931, 1989.

[3] Hui Li, B.S.Manjunath and Sanjit K. Mitra, "MULTI-SENSOR IMAGE FUSION USING THE WAVELET TRANSFORM," in *Proceedings of ICIP '94*, PP. 51-55, 1994.

[4] N. Nandhakumar and J. K. Aggarwal, "Integrated Analysis of Thermal and Visual Images for Scene Interpretation," in *IEEE Trans. on PAMI*, Vol. 10, No. 4, pp. 469-481, July 1988.

[5] Axel Pinz and Manfred Prantl, "Active Fusion for Remote Sensing Image Understanding," in *SPIE Proceedings of Image and Signal Processing for Remote Sensing II*, Paris, 1995, Vol. 2579, pp.67-77, 1995.

[6] Veronique Clement, Gerard Giraudon, Stephane Houzelle, Fadi Sandakly, "Interpretation of Remotely Sensed Images in a Context of Multisensor Fusion Using a Multispecialist Architecture," in *IEEE Trans. on Geoscience and Remote Sensing*, Vol. 31, No. 4, July 1993.

[7] B. Chadrsekaran, "Generic Tasks in Knowledge-Based Reasoning : Characterizing and Designing Expert Systems at the Right Level of Abstraction," in *Proceedings of Conf. AI and Applications*, 1985, pp. 294-299

[8] B. Chadrsekaran, M. C. Tanner and J. R. Josephson, "Explaining Control Strategies in Problem Solving," *IEEE Computer*, Spr. 1989, pp. 9-24