
PROCESAMIENTO DEL LENGUAJE NATURAL

BOLETIN N.º 1

OCTUBRE 1983



EDITA Y DISTRIBUYE: FACULTAD DE INFORMATICA

APDO. 649

SAN SEBASTIAN

PROCESAMIENTO DEL LENGUAJE

BOLETIN Nº 1

OCTUBRE 1.983

EDITORIAL	1
PRESENTACION DE LINEAS DE TRABAJO	
Un método para la eliminación de la redundancia en la transmisión de textos escritos en castellano.	5
F. Ares de Blas. Facultad de Informática. San Sebastián	
Esquema - resumen	10
N. Antonio Campos. Colegio Universitario. Ciudad Real	
Procesamiento del habla	12
F. Casacuberta y E. Vidal. Centro de Informática. Valencia	
Análisis morfológico automático del español	18
Montserrat Meya. Siemens. Munich	
Prospecció Automatizada de Textos Catalans	24
J. Rafel. Dpto. Catalán. Universidad de Barcelona.	
Etiquetado no fonemático para un sistema automático monolocator de reconocimiento de palabras aisladas.	33
J. Rubio Ayuso y M.C. Carrión Pérez. Universidad de Granada.	
Esquemas - Resumen	37
Martín S. Ruipérez. Facultad de Filología. Universidad Complutense	
USL: Un sistema para interrogar en castellano a bases de datos relacionales.	38
L. de Sopeña. Centro Científico IBM. Madrid	
Q - Séneca. Un sistema pregunta-respuesta para el castellano.	42
M.F. Verdejo. Facultad de Informática. San Sebastián.	

INFORMACION

Acta de la reunión de Sevilla. Abril 1.983	47
Personas. Centros de Trabajo. Recursos.	51
Estatutos de la Sociedad Española para el procesamiento del lenguaje natural.	52
Noticias de Congresos	54
Sección Bibliográfica	66

EDITADO POR: M. MEYA y M. F. VERDEJO

COMPOSICION Y DISTRIBUCION: FACULTAD DE INFORMATICA

Apartado 649

SAN SEBASTIAN

EDITORIAL

La aparición de este primer número de la revista Procesamiento del Lenguaje se encuentra estrechamente ligada al nacimiento y constitución de la Sociedad Española para el Procesamiento del Lenguaje Natural.

El propósito que perseguimos con ella es, ante todo, crear un órgano especializado de comunicación a través del cual podamos conocer de forma directa los trabajos, investigaciones y proyectos que se realizan en España dentro del área del Procesamiento de lenguaje Natural. Por ello, al tratarse de un tarea interdisciplinaria, nos hemos reunido en un mismo empeño profesionales de diferente formación: ingenieros, informáticos y lingüistas.

Queremos que el lenguaje sea tratado en toda su complejidad:

- a) como signo, es decir, como entidad física bien sea óptica (gráfica) o acústica.
- b) como proceso, considerando todo tipo de actividad intelectual : procesos memorísticos, cognitivos, asociativos, etc...
- c) como almacenamiento estático del conocimiento; incluye todo tipo de representación formal : redes semánticas, frames, representación en lógica formal, etc...
- d) como medio ya sea como instrumento de manipulación, de persecución de objetivos, ya sea considerando la situación conversacional, pragmática, etc...

En todos estos casos nos ocupamos del tratamiento del lenguaje mediante métodos algorítmicos, y procesado con ayuda de ordenadores electrónicos.

Si consideramos el lenguaje como entidad física sonora (señal acústica) nos ocuparemos del lenguaje oral hablado o habla. Esto incluye diferentes campos de investigación: reconocimiento del habla (palabras aisladas, frases o conversaciones) identificación de parlantes, codificación (parametrización), síntesis del habla, ... Incluye, pues, también el paso automático de uno a otro tipo de lenguaje, oral a escrito y viceversa.

Como entidad física gráfica - el lenguaje escrito - el lenguaje puede ser tratado a nivel de cadena, es decir, se usan criterios y medios que nos ayuden a vaciar textos, catalogarlos, describirlos, interpretarlos, criticarlos, especificar la autoría, etc... El resultado de trabajos de este tipo serán: índices, concordancias, listas comparativas, obtención de material contrastivo, tablas, frecuencias, resultados estadísticos, especificación de sublenguajes, vocabulario básico, etc...

Al considerar el lenguaje como proceso se trata de especificar un conjunto de comportamientos a partir de reglas bien definidas, así como definir funciones que transformen procesos en otros procesos. Igualmente entran en consideración el aprendizaje, el procesamiento de primitivos del conocimiento, su reflejo e interacción con el saber lingüístico, estrategias de inferencia para comprender un texto, etc...

El tratamiento del lenguaje como representación del conocimiento en el que interaccionan saber

lingüístico, situacional, enciclopédico, etc,... implica el diseñar modelos que de forma aproximada reflejen, ya la organización del conocimiento en el ser humano, ya la estructura conceptual de un texto o una conversación.

Por último, si se procesa el lenguaje considerándolo como medio esto implicará tomar en consideración aspectos como: pragmática situacional, objetivos que persigue el hablante (sin los cuales no se comprende el texto o diálogo), etc...

Esperamos, por tanto, dar cabida en nuestra revista a un amplio espectro de actividades que, no obstante, a partir de un planteamiento teórico busquen una orientación y aplicación práctica a un ordenador.

Este primer número es básicamente una exposición sucinta de los trabajos que actualmente realizamos unos cuantos miembros del grupo, que en Abril se reunió en Sevilla en la cátedra de Lengua Española.

En el futuro confiamos que la presente revista no sirva exclusivamente para dar a conocer las actividades y proyectos que en el área del Procesamiento del Lenguaje Natural se llevan a cabo en nuestro país, sino también para dar noticia de los trabajos que se realizan en este campo en el extranjero. Se incluirá información sobre congresos internacionales, simposios, mesas redondas y grupos de trabajo, incluyendo además reseñas sobre libros y artículos de especial interés.

Nuestra mayor satisfacción será saber que nuestra revista sirve de motivación para los diferentes grupos y personas que trabajaban hasta ahora totalmente aislados, y con ella poder activar la

colaboración, comunicación e intercambio entre todos ellos.

Como iniciadora de aquella reunión, y ahora de esta publicación agradezco el apoyo del Dr. V. Lamiquiz (Univ. de Sevilla), de M.A. Pineda, de Jordi Romano, la colaboración de aquellos que incluyen aquí sus exposiciones, y sobre todo la de la Dra. F. Verdejo con quien a tandem hemos hecho la redacción.

M. Meya

Un método para la eliminación de la redundancia en la transmisión de textos escritos en castellano

Por Félix Ares de Blas

A lo largo del presente trabajo se han conseguido algunas modestas aportaciones en el campo de la eliminación de la redundancia en transmisión de textos escritos en castellano. Las podríamos resumir así:

- 1º.- Hemos calculado unos límites superior e inferior a la cantidad de información que proporciona una letra del castellano impreso, siguiendo el método que utilizó Shannon para el idioma inglés. En principio dicho conocimiento permite a los investigadores de este tema saber dónde está su límite teórico.
- 2º.- Hemos constatado que refiriéndonos exclusivamente a la cantidad de información de una letra o de un conjunto de letras, el inglés y el castellano tienen valores muy similares, por lo que los datos existentes para el idioma inglés pueden extrapolarse, con las debidas precauciones al castellano.
- 3º.- Hemos constatado, aunque a un nivel todavía sujeto a discusiones, que cada letra de un texto traducido da menos información que la del texto original.
- 4º.- Hemos demostrado que hay unas limitaciones de orden teórico a la codificación de la extensión de orden n de una fuente. n no puede crecer indefinidamente. Hemos probado que, dependiendo de la longitud del texto a codificar y de los diccionarios utilizados en su codificación, hay un n óptimo. Con ello hemos clarificado la gran influencia que tiene la inestabilidad de las frecuencias relativas de aparición de n - gramas.
- 5º.- Hemos constatado que la transmisión de textos rompiéndolos en palabras sólo es adecuado para fuentes y diccionarios muy grandes, que para fuentes o diccionarios de pequeño tamaño, codificar a nivel de palabra puede ser peor que codificar a nivel de 3 - gramas.

- 6º.- Hemos visto que la sílaba es un buen bloque para conseguir un ahorro de binitis en transmisión de textos escritos en castellano. Se trata de un bloque grafémico multilettra frecuente y estable. Como consecuencia del estudio silábico - hemos desarrollado unos programas que rompen las palabras en sílabas.
- 7º.- Hemos visto que la información que proporciona una letra - depende de su posición dentro de la palabra, y hemos desarrollado algunos métodos para obtener una ventaja de este hecho.
- 8º.- Hemos demostrado que la información que proporciona una sílaba depende de su posición dentro de la palabra, y hemos desarrollado algunos métodos para obtener un cierto ahorro de binitis a partir de este hecho. Concretamente el método de distinguir entre tres tipos de sílabas posicionales, su poniendo que el blanco de separación entre palabras forma parte de la sílaba final, ha demostrado ser el mejor para casi todos los casos estudiados, llegándose a los 2'91 --- binitis/letra. Solamente ha sido superado en el caso de textos fuente y diccionarios muy grandes por la codificación en palabras.
- 9º.- Como conclusión podemos decir que hemos desarrollado un método de eliminación de la redundancia para textos escritos en castellano totalmente original, muy adaptado a dicho -- idioma, y con el que se consiguen resultados similares, ligeramente mejores que los reportados para el inglés. (Codificado y transmitido: 2'91 binitis por letra. Teniendo en cuenta sólo la redundancia, sin llegar a codificar: 2'61 - bits de información por letra).
- 10º.- Marginalmente, como subproductos de la investigación, hemos obtenido ciertos datos estadísticos sobre textos escritos en castellano. Por ejemplo, frecuencias relativas --- (f.r.) de letras, f.r. de letras según su posición en la - palabra tanto empezando a numerar por la izquierda como -- por la derecha, palabras diferentes que constituyen el 25% y el 50 % de un texto, f.r. de palabras, f.r. de sílabas, f.r. de las sílabas de acuerdo con su posición dentro de - la palabra, f.r. de las letras dentro de las sílabas te---

niendo en cuenta su posición en ellas, f.r. de las letras dentro de las sílabas teniendo en cuenta su posición en -- ellas y considerando tres tipos diferentes de sílabas según su posición dentro de la palabra, sílabas que forman - 25 % y el 50 % de un texto separando por los diferentes tipos posicionales de sílabas, sílabas más frecuentes teniendo en cuenta sus diversos tipos, cálculo de las entropías de todos los casos en que había f.r., longitudes medias de palabras, de sílabas, de sílabas posicionales, f.r. de aparición de letras dentro de n-gramas (Para $n = 3, 4, 5$ y 6), f.r. de n-gramas, costos relativos de transmisión en binitis (dígitos binarios) y en bits (unidades de información) de n-gramas, palabras, sílabas con el blanco considerado sílaba aparte, sílabas con el blanco integrado en la última sílaba, sílabas posicionales con el blanco como última sílaba, sílabas posicionales con el blanco integrado en la última sílaba, costo relativo de inestabilidad de los mismos casos que los ya señalados para costos relativos de transmisión. Además, frecuencias relativas de palabras de n letras, frecuencias relativas de sílabas de n letras, f.r. de palabras de n sílabas. Codificación Huffman de diversos alfabetos castellanos (con 27 letras, con 28, etc.), codificación de Huffman de las longitudes de las palabras medidas en letras y en sílabas, etc., etc.

TEXTOS BASE USADOS

Para realizar la investigación hemos partido de la -- elección aleatoria de una serie de autores, para los que hemos perforado parte de su obra. La perforación se ha hecho en código Fieldata, sólo se han tenido en cuenta las mayúsculas, se -- tienen en cuenta varios signos de puntuación que se intercalan entre las palabras: punto, punto y coma, punto y aparte, dos puntos, abrir admiración, cerrar admiración, abrir y cerrar admiración, paréntesis, etc., pero no se han tenido en cuenta los signos de puntuación que van sobre las letras: acentos, diéresis, etc.

Las obras han sido (en el orden de elección aleatorio):

- José M^a Gironella.- "Un millón de muertos". Editorial Planeta. Colección Omnibus. Barcelona 1966. 139 Edición. (Cap. 1).
- Ramón Menéndez Pidal.- "Los españoles en la historia". Ed. Espasa Calpe (Argentina) S.A. Colección Austral. (Vol. N^o 1.260). Buenos Aires 1959. (1^a Edición). (Cap. 1)
- Pío Baroja.- "La Trapera" y la "La sima" Ed. Alianza Editorial. Colección el libro de bolsillo. Madrid 1967. 2^a edición (Primeras 12.628 letras que comprenden íntegramente "La trapera" y parte de "La sima".
- Miguel de Unamuno.- "Amor y Pedagogía". Ed. Emesa. Colección Noveles y Cuentos. Madrid 1967. 2^a edición. (Primeras 8.426 letras del capítulo 2).
- Alvaro de Laiglesia.- "Sólo se mueren los tontos". Ed. Planeta. Colección novela. Barcelona 1967. (10^a Edición). (Pedazo IV).
- Benito Pérez Galdós.- "Marianela". Ed. Hernando. S.A. Colección Novelas españolas contemporáneas. 1^a época Madrid 1943. (Capítulo 1)=

Camilo José Cela.- "La familia de Pascual Duarte".

Ed. Destino. Colección: Destinolibro. Vol 4.
Barcelona 1977. (6ª edición). (Capítulo 3)

Además añadimos un texto periodístico al azar. La ---
elección recayó sobre el País. Editorial "Voto de conciencia" -
de Marzo de 1981 y el artículo "Prudencia y jurisprudencia" del
15/7/81.

Para tener una referencia clásica utilizamos "Don Qui
jote de La Mancha" de Miguel de Cervantes. Editorial Bruguera.
Libro clásico. Barcelona 1981. La obra consta de dos volúmenes
se perforó el primero de ellos completo, incluyendo el prólogo
de Juan Alcina Franch.

ESQUEMA - RESUMEN

NICOLAS ANTONIO CAMPOS

Departamento de Lingüística Francesa

Colegio Universitario (Universidad Complutense)

Ciudad Real

Siendo consciente de mis limitaciones, ya que por una parte, mi formación lexicométrica es autodidacta y mis conocimientos en informática son pobres, y por otra, el tener que desplazarme a Madrid (Centro de Cálculo de la Universidad Complutense) para tener acceso a los medios técnicos necesarios para llevar adelante las investigaciones me obliga a ser muy modesto en cuanto a los objetivos.

Trato de realizar el estudio del vocabulario político de 200 textos (250.000 ocurrencias) de 5 publicaciones francesas de periodicidad semanal y de ámbito estatal que durante 1.968 aporten, divulguen o comenten la revolución estudiantil de mayo.

Una opción fundamental guía nuestra práctica: El texto como dato debe ser, en la medida de lo posible, fiel al documento sobre el cual trabajamos. Esto implica:

- a) Un registro exhaustivo
- b) Una precodificación reducida
- c) Una asepsia controlada

. Clasificar las formas del texto por orden alfabético indicando su frecuencia absoluta y sus referencias.

. Clasificar las mismas formas por orden de frecuencia decreciente, haciendo figurar su rango y su frecuencia relativa.

Ajustarnos rigurosamente a este punto de vista, retener todo y de la misma manera será imitar el comportamiento de un paleógrafo intentando descodificar un texto en lengua desconocida.

Salvando algunas dificultades para desambiguar el texto y procurando dar un estatuto general a los signos procederemos - de la manera siguiente:

1. Automatización del texto

- 1.1 Registro del texto
- 1.2 Registro de información complementaria al texto, fichero de parámetros.
- 1.3 Indexación
- 1.4 Delimitación y fragmentación del corpus
- 1.5 Resultado de la indexación. Corrección.

2. Estadística de frecuencias

- 2.1 Índice alfabético de formas léxicas y funcionales
- 2.2 Índice jerárquico por orden de frecuencia decreciente con indicación de F.A. y F.R.
- 2.3 Estudio de la variabilidad de frecuencia de las formas. Test apropiados.
- 2.4 Índice de reparto
- 2.5 Formas de base y formas específicas (+ y -)

3. Localización

- 3.1 Las coocurrencias y los contextos
- 3.2 Relación de las "Couples"
- 3.3 Relación de las "Couples" en contexto
- 3.4 Parejas coocurrentes

Análisis por campos, comparación de subcorpus. Análisis de las formas léxicas más frecuentes y los hapax. Ausencia/presencia.

PROCESAMIENTO DEL HABLA
Francisco Casacuberta y Enrique Vidal
CENTRO DE INFORMATICA DE LA UNIVERSIDAD DE VALENCIA

1.- INTRODUCCION.

El reconocimiento de la palabra continua es una de las facultades que caracteriza la inteligencia humana y en la actualidad se es consciente de la enorme complejidad que supone la concepción de sistemas que intenten aproximarse a las prestaciones del cerebro humano. La investigación a desarrollar en este dominio es, pues, a la vez fundamental (estudios de las propiedades y mecanismos del habla) y aplicada (desarrollo de sistemas prácticos).

En la figura adjunta se recogen esquemáticamente los proyectos en curso de desarrollo en el C.I.U.V., junto con los problemas que ellos abordan.

2.- LINEA FUNDAMENTAL.

En el aspecto informático de la rama fundamental, se abordan problemas de representación del conocimiento y se trabaja en el desarrollo de técnicas de Reconocimiento Sintáctico de Formas y en la aplicación de la Teoría de Conjuntos Difusos a la descripción sintáctica de las formas acústicas del habla; mientras que en el aspecto de Tratamiento de la Señal se investigan nuevos modos de caracterizar los segmentos transitorios de la señal vocal.

2.1.- RECONOCIMIENTO SINTACTICO.

En este punto se estudia, en primer lugar un sistema que permita una descripción sintáctica-difusa de los elementos fonéticos o subfonéticos del habla, esto es, partiendo de una familia indexada de conjuntos de parámetros (segmentos), extraídos mediante ventanas que se deslizan en el tiempo, se trata de obtener otra familia indexada de subconjuntos difusos de un determinado conjunto de símbolos fonéticos o subfonéticos (Etiquetado Difuso), mediante un descriptor y un traductor difuso (1).

Con el etiquetado difuso se cambia la representación paramétrica clásica por una representación simbólica de la señal vocal.

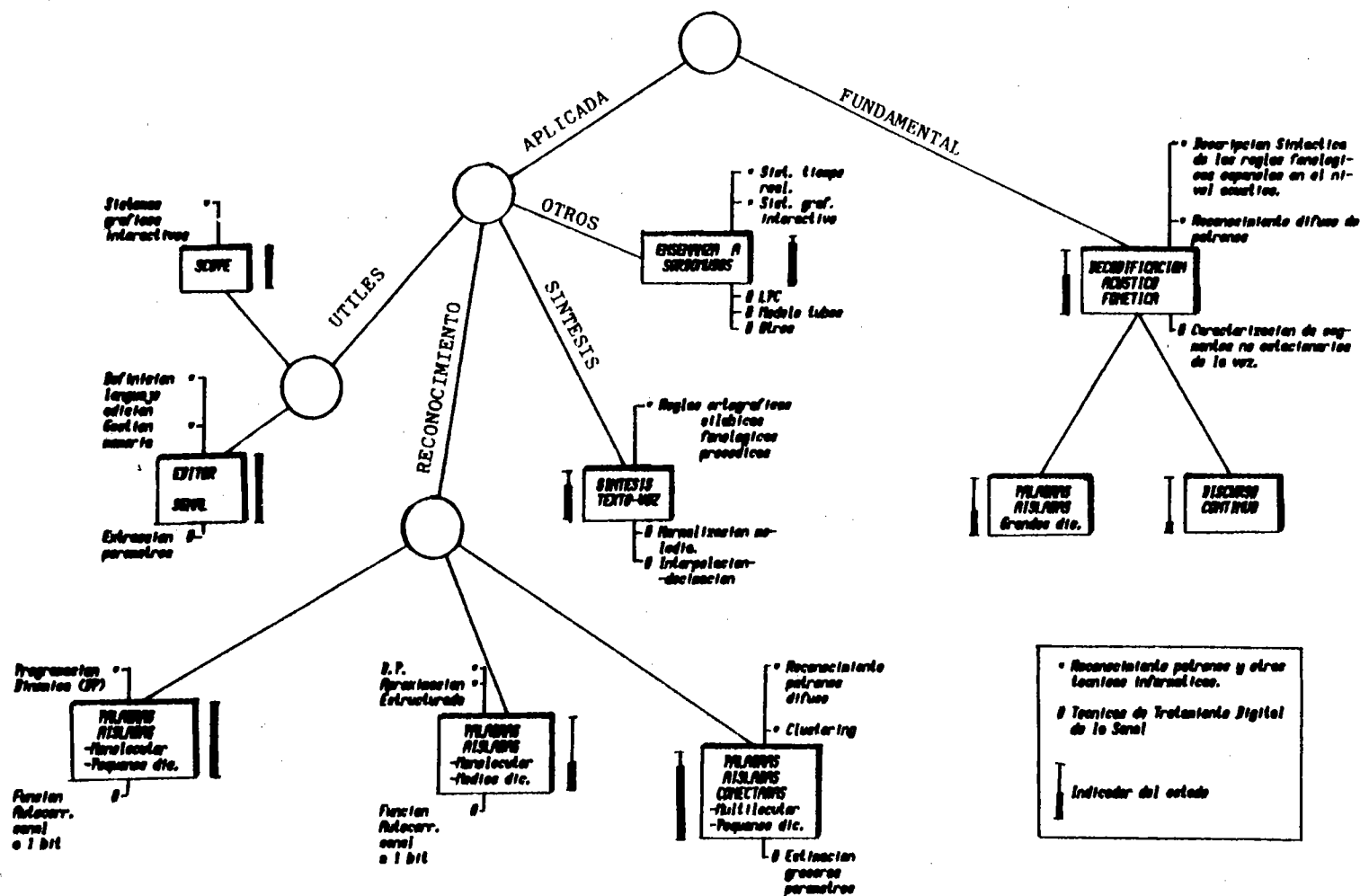
En segundo lugar se estudian los sistemas de reconocimiento propiamente dichos, esto es, partiendo de una descripción sintáctico-difusa de una palabra pronunciada y de una descripción fonológica, como representación del conocimiento, se obtiene un subconjunto difuso de palabras que forman el diccionario, donde la palabra de mayor evidencia puede ser considerada como la palabra más próxima a la pronunciada según el sistema (2).

Estos trabajos han sido aplicados a pequeños sistemas experimentales que serán descritos más adelante.

2.2.- CARACTERIZACION DE SEGMENTOS NO ESTACIONARIOS DE LA VOZ.

Uno de los problemas angulares en la decodificación acústico-fonética es el reconocimiento de los segmentos explosivos y nasales. Dichos segmentos son de gran importancia debido a su gran número en la mayoría de las lenguas, su caracterización resulta de extremada complejidad debido a su carácter no estacionario, y normalmente las soluciones que se aportan no se dan en el nivel

ESTADO DE PROYECTOS Y ACTIVIDADES DE INVESTIGACION Y RECONOCIMIENTO EN EL C.I.U.U. (1983)



acústico.

El método que se propone se basa en los estudios de la evolución en el tiempo de los parámetros de los formantes que caracterizan a las vocales siguientes al sonido consonántico (3), y en un modelo para la detección de estos parámetros en segmentos estacionarios de voz, (4), y en el cual suponemos una variación lineal de éstos; así pues, en nuestro caso, un formante, en un segmento no estacionario, viene definido por la amplitud y frecuencia central del polo y por sus correspondientes derivadas temporales. Para la obtención de dichos parámetros se han propuesto dos modelos: el primero se basa en una variación apreciable de período a período, realizándose un análisis frecuencial con dos períodos de voz (5) mientras que en el segundo se suponen fuertes variaciones dentro de un período, aplicándose un análisis frecuencial continuo en un período (6).

Aunque los resultados obtenidos son parciales, éstos hacen suponer que los modelos escogidos darán buenos resultados en un futuro próximo.

3.- LINEA APLICADA.

En esta línea se abordan proyectos prácticos en los cuales tienen aplicación inmediata los resultados de las investigaciones fundamentales. Caben destacar dos sistemas ya en funcionamiento, de Reconocimiento de Palabras Aisladas, uno monolocator y otro multilocator para diccionarios específicos. La eficacia y simplicidad de estos nuevos sistemas, permite su implementación sobre microprocesadores de propósito general, sin prácticamente ningún material adicional.

Otros proyectos abordados en esta rama abarcan temas desde la Edición Interactiva de señales vocales, hasta la Enseñanza del Habla a sordomudos, pasando por un Sistema de Síntesis sin restricciones de frases pertenecientes a cualquier texto castellano ortográficamente correcto.

3.1.- RECONOCIMIENTO DE PALABRAS AISLADAS (R.P.A.).

3.1.1.- SISTEMA DE R.P.A. MONOLOCUTOR.

Este proyecto parte de un estudio experimental realizado en nuestro Laboratorio, en el cual se demuestra que la función de autocorrelación de la señal cuantizada a 1 bit, puede ser utilizada como método de parametrización de un Sistema de Reconocimiento de Palabras Aisladas (7). Esta función es de una complejidad de cálculo pequeña y puede ser implementada con un material de bajo coste o bien en ensamblador, permitiendo obtenerla en tiempo real en un (micro) miniordenador de propósito general.

Para la fase de decisión se implementó un algoritmo de Programación Dinámica simple y rápido: el de forma simétrica sin restricciones de pendiente y con peso unidad en el camino diagonal (8).

Las tasas de reconocimiento obtenidas de este sistema son del orden de 97% para diccionarios de unas 30 palabras (9).

En este momento se está abordando el proyecto de implementar un sistema real (un mini-query por la voz), basado en el sistema aquí expuesto brevemente.

3.1.2.- SISTEMA DE R.P.A. MULTILOCUTOR BASADO EN EL ETIQUETADO DIFUSO Y EN PROGRAMACION DINAMICA.

El etiquetado difuso, estudiado en el apartado 2, utiliza, en este caso, como parámetros a la amplitud media y a la densidad de cruces por cero de un segmento de voz, y como etiquetas un conjunto de categorías fonéticas "groseras".

El tipo de representación indicado, permite, en primer lugar, la obtención del diccionario mediante traducción de las cadenas ortográficas en cadenas de símbolos que representan categorías fonéticas "groseras", utilizando las reglas fonológicas correspondientes. Y en segundo lugar, una representación contigua de aquellas cadenas que comiencen con las mismas subcadenas (9).

En la etapa de reconocimiento, propiamente dicha, se ha utilizado un algoritmo de Programación Dinámica, que aprovecha la simplicidad de dicho tipo de representación.

Las tasas de reconocimiento obtenidas con este sistema son bajas, alrededor del 70 %. Para un sistema real se pueden mejorar los resultados simplemente aumentando el poder descriptivo de la representación mediante el uso de parámetros complementarios.

3.1.3.- SISTEMA DE R.P.A. MULTILOCUTOR BASADO EN METODOS SINTACTICOS

El sistema desarrollado en este proyecto está basado en el etiquetado difuso citado en el apartado anterior, y en métodos difuso-sintácticos como sistema de decisión, utilizando una descripción sintáctica del léxico basada en las categorías acústico-fonéticas mencionadas, al cual se le asocia un algoritmo de comparación de la cadena de etiquetas difusas de entrada con la representación de las palabras del léxico. Una descripción formal de este sistema puede encontrarse en (9).

El sistema de reconocimiento descrito en este apartado ha sido implementado con una serie de modificaciones y experimentado para tres diccionarios: dígitos castellanos, dígitos catalanes y un lenguaje mini-query, y han sido obtenidas unas tasas de reconocimiento, del 85 %, resultado elevado teniendo en cuenta la "pobreza" de los parámetros utilizados.

Se han utilizado los mismos parámetros que en el sistema anterior para la obtención del etiquetado difuso, habiéndose previsto la incorporación de un tercer parámetro (densidad de cruces por cero de la señal derivada) que permitirá subdividir las categorías fonética utilizadas, y por tanto una mejor descripción del léxico.

3.2.- SINTESIS AUTOMATICA DEL CASTELLANO HABLADO.

Para contemplar el proceso inverso al Reconocimiento del Habla, se está desarrollando un proyecto de síntesis automática a partir de un texto escrito. Este proyecto se basa en decodificación de una cadena constituida por determinados elementos discretos (difonemas) concatenados y a los cuales se les ha aplicado determinadas reglas para obtener una mayor similitud al castellano hablado. El sistema utiliza un diccionario de difonemas

que consta de señales en el dominio del tiempo digitalizado (10).

El sistema se compone de una serie de etapas, algunas de las cuales se realizan en paralelo; así, a partir del texto escrito, se realiza una división silábica obteniéndose a continuación una transcripción fonética y de esta, a su vez, una transcripción difonética, con la cual se puede realizar la extracción del conjunto de difonemas del diccionario. En paralelo con las etapas descritas se obtiene a partir del texto, la información prosódica (el acento y la entonación) que una vez procesada intervienen en la concatenación de los difonemas y las variaciones de frecuencia de salida de muestras del conversor digital analógico. En la actualidad este último conjunto de etapas aun no ha sido implementado.

3.3.- AYUDA A LA ENSEÑANZA DEL HABLA A SORDOMUDOS.

Una de las aplicaciones de mayor interés social en este área de investigación, es el diseño de sistemas de ayuda a minusválidos y, en particular para el aprendizaje del habla de sordomudos.

La idea general es la de suplir la realimentación que se realiza mediante el sistema auditivo por un gráfico en el cual se muestra de alguna forma el resultado de la pronunciación de un sonido por un locutor. En nuestro caso mediante la representación en tiempo real, de un modelo de 8 tubos del aparato fonador humano, que es representado en el terminal de ordenador, teniendo a la vez la de un patrón con el mismo modelo.

El sistema desarrollado en nuestro Laboratorio está en fase de prueba con diferentes individuos sordomudos.

3.4.- UTILES

3.4.1.- SISTEMA GRAFICO INTERACTIVO: SCOPE.

SCOPE es un sistema gráfico interactivo orientado específicamente a la representación por una pantalla de rayos catódicos de señales y-z-t, y está basado en un microprocesador con un material específico para controlar la citada pantalla de bajas prestaciones. Las principales características son: flexibilidad de uso, fácil instalación en un terminal alfanumérico asíncrono y un bajo coste de implementación (11). Los dos subsistemas que componen SCOPE son: Un sistema básico y uno auxiliar. El primero, basado en el microprocesador, puede ser considerado como un dispositivo insertado en la línea asíncrona que une el terminal alfanumérico y un ordenador. El sistema auxiliar reside en el ordenador principal y permite el diálogo con el sistema base, pudiendo obtenerse copias permanentes en fichero o sobre trazador X-Y ("plotter") (11).

3.4.2.- EDITOR DE SEÑAL (EDS).

El Editor de Señal desarrollado en nuestros laboratorios es un útil de edición de señales cualesquiera, y en particular para el tratamiento de señales acústicas. Las principales aplicaciones del EDS son: creación de diccionarios de elementos acústicos y verificación de resultados en Síntesis, segmentación y marcaje manual en Reconocimiento, etc.

El EDS lleva asociado por un lado un Leguaje de Comandos que permite describir tratamientos para la manipulación de la señal, tal como inserción,

borrado, inversión, etc., así como un comando general definido por el usuario, para permitir cualquier manipulación de la señal. Los objetos del Lenguaje de Edición son representaciones de la señal acústica que admiten clasificaciones en diversos "tipos" que son referenciadas por nombres simbólicos (variables)(12).

4.- CONCLUSIONES.

Algunos de los trabajos presentados pueden darse por finalizados, sirviendo en este momento para el diseño de maquetas reales, o como herramienta de investigación. Otros en cambio están en fase de estudio y/o desarrollo, y su aplicación será a medio plazo. Todos ellos son el fruto del esfuerzo de un grupo de 6 personas que desde finales de 1980 han dedicado parte de su tiempo a la investigación e implementación en el campo del Reconocimiento y Síntesis Automática del Habla, construyendo los temas de varias tesis realizadas o en curso de elaboración.

BIBLIOGRAFIA.

- (1) VIDAL E., CASACUBERTA, F., RULOT H., SANCHIS E. " Isolated Word Recognition based on Fuzzy Labelling and Dynamic Programming" 1983 SWSSPA. Sept 83 Sitges.
- (2) VIDAL E., SANCHIS E., CASACUBERTA F. " A Speaker-Independent Isolated Word Recognition System for Specific Dictionaries" 1983 SWSSPA. Sep 83. Sitges.
- (3) RUSKE G. "Auditory Perception and its Application to Computer Analysis of Speech" Computer Analysis and Perception. Vol II. 1981.
- (4) FRIEDMAN D. "Estimation of Formant Parameters from a Sum of Pole Modeling" Proc. ICASSP-81, pp. 351-354. Apr-81.
- (5) VIDAL E., CASACUBERTA F. " Incremental Linear Model for Non Stationary Voiced Speech Segments" 11.ICA.Toulouse. Jul. 83.
- (6) CASACUBERTA F., VIDAL E. "A Continuous Linear Model for Non Stationary Analysis of Voiced Speech Transitions: an Advanced". 1983 SWSSPA. Sep. 83 Sitges.
- (7) RULOT H., VIDAL E., CASACUBERTA F. "La Función de Autocorrelación en el Reconocimiento de la Palabra" V Congreso de Inf. y Aut. Madrid 81, pp. 799-803.
- (8) RULOT H., VIDAL E., CASACUBERTA F. " Isolated Word Recognition System based the Autocorrelation Function ". 1982 PWSSPA. Oct. 82 Pavia de Varzim.
- (9) VIDAL E., CASACUBERTA F., SANCHIS E., RULOT H. " Diversas Aproximaciones al Reconocimiento de Palabras Aisladas ". V Escuela de Verano de Informática Cáceres. Jul. 83.
- (10) TORRES B., VIDAL E. " Sistema de Síntesis automática del Castellano Hablado " V Congreso de Inf. y Aut.. Madrid 81.
- (11) VIDAL E., TORRES B., CASACUBERTA F. " Un Sistema Gráfico Interactivo para la Representación de Señales Unidimensionales: SCOPE " Rev. de Informática y Automática. Pendiente de publicación.
- (12) BENEDI J. " Un Editor de Señal ". Tesina de Licenciatura. Mar. 1983. Univ. Valencia.

Analisis morfologico automatico del espanol
Montserrat Meya
Munich

- 1 Objetivos
- 2 base de datos
- 3 componentes

1 Objetivos

El análisis morfológico de las palabras de un texto puede realizarse:

- a) a nivel de cadena, las palabras son despojadas de 1,2,3,...n grafemas mediante algoritmos que asignan categorías, sin diccionario.
- b) a nivel de palabra, bien con diccionario de formas de palabras, o bien con diccionario de lemas- formas base - . En este caso se trata fundamentalmente de búsqueda de diccionario reduciéndose el número de algoritmos.
- c) a nivel de morfema, con un diccionario mucho más reducido - hay muchos menos morfemas que palabras- pero con mayor aparato procedural.

Un análisis morfológico a nivel de morfema -tipo c- tiene la ventaja de:

- necesitar menos memoria para el almacenamiento de los datos lingüísticos -menor número de entradas-.
- poder relacionar entre sí diversas formas de palabras con una raíz común, aun en aquellos casos en que se trata de formas fuertes o irregulares.

El análisis morfológico aquí expuesto puede ofrecer:

- la descomposición de cada palabra de un texto escrito en sus morfemas respectivos
- la categorización morfológica
- lematización o reducción de las formas de palabras a su forma base(p.ej.: "trabajaba" ---->"trabajar")
- especificación de la raíz (p.ej.: "trabajaba"--->"trabaj")

El objetivo del analizador morfológico es obtener descriptores que den cuenta del tipo de información tanto de los documentos-textos almacenados- como de la demanda formulada por el usuario a un banco de datos.

El "recall" es obviamente enorme, si no se recurren a otras funciones de retrieval que precisen o ajusten la información a considerar.

2 Base de datos

El modelo de análisis morfológico automático es aplicable a cualquier lengua -está implemantado para el alemán, español e inglés- sólo la base de datos y pocos programas específicos varían según la lengua de que se trate.

La base de datos consta:

- A : Listado de morfemas categorizados a partir de propiedades
 - léxico-sintácticas (raíz léxica o no, nominal, verbal,...)
 - flexivas : tipo de flexión, tipo de variante,...
 - distributivas: reglas de serialización, posición en cadena..)
- B : tabla de condiciones para las transformaciones(cuento---> contar; transformaciones ue --> o, ie-->e, g-->j, etc...)
- C : Tabla de formas fuertes (alomorfos) a los que se hace referencia por un puntero.
- D : Tabla de formantes flexivos con información detallada a usar en el posterior análisis sintáctico.
- E : Palabras funcionales (gramaticales)- stop-words.
- F : Tabla de reglas para la descomposición.

2.1 Especificación subcategorización de los datos

En la cadena gráfica de un asiento pueden coincidir varios morfemas. Cada asiento es ,pues, un morfo. Por ej.

"libr" representa 4 morfemas : libro.libra.libre.librar

Cada entrada/asiento tiene asignado el siguiente patrón de 32 bits.

LEX	NOM	VRB	ADJ	ADV	NUM	FUN	GEN	Conj.	C	Trafo		ACT	D-0	
ø	o	a	e	o/a	ø/a	ALO	DIF	AFF	PRE	ADn	FLX	SUF	Subcat	Adv

Categoriz.morfológica: tipo de morfema: nominal,verbal,etc,
Subcategorizacion: GENero,Conjugacion(bits 9-10),tipo de flexión
(bits 17-22),alomorfo,transformaciones posibles,SUFijo,Adnnomi-
nal,etc...

Los alomorfemas que por procesamiento pueden ser reducidos a su morfema originario no tienen asiento. Por ejemplo tiene asiento "conoc"-er, pero no "conozc-o" ni tampoco "cuent"-o que son transformados en el proceso de reconocimiento.

Ejemplo de asientos del diccionario:

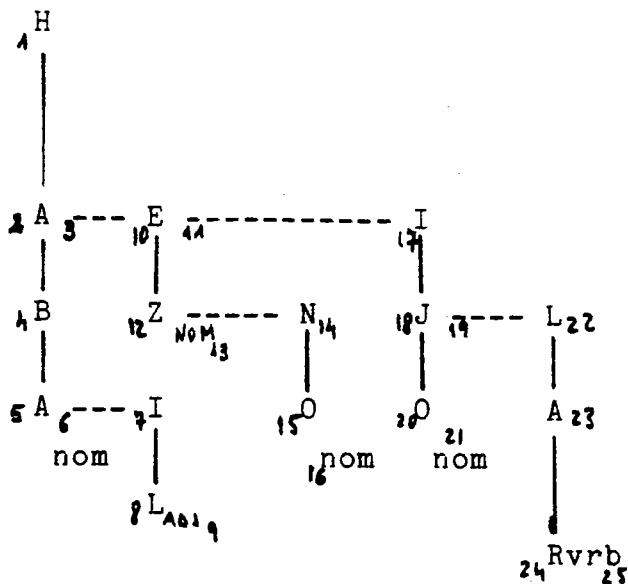
ventan : LEX NOM a
cont : LEX NOM VRB o a ar ue D-0
alguien: FUN
o : FUN AFF FLX NOM VRB ADJ ADV NUM o o/a ar er ir
ción : AFF SUF NOM 0 Deverbal,....

La tabla de transformaciones es procesada por SPATZ. Esta tabla tiene almacenada en forma arborea las subcadenas que permiten transformacion junto con una categoria (p.ej:1=c-->z)

2.2. Almacenamiento de los datos

Las diversas listas estan comprimidas/almacenadas en forma arborescente lo que ademas facilita la búsqueda. Cada letra del alfabeto tiene un árbol que será más o menos extenso según el número de ramificaciones que tenga. El árbol para "p" es mayor que el de "h" porque hay muchos más morfemas que empiecen por tal letra. Se trata de una optimación de una representación arborea bisecuencial. Las siguientes palabras estarían almacenadas:

haba
 habil
 hez
 heno
 hilar
 hija



1	H	↓
2	A	↓
3		10
4	B	↓
5	A	→ ■
6	NOM	
7	I	↓
8	L	→ ■
9	ADJ	
10	E	↓ ►
11		17
12	Z	→ ■
13	NOM	
14	N	↓
15	O	↓ ■

2.3 Tabla de reglas para la descomposicion

La descomposición de una palabra en sus morfemas se hace a partir de 15 posibles estados y las posibles transiciones entre estos estados:

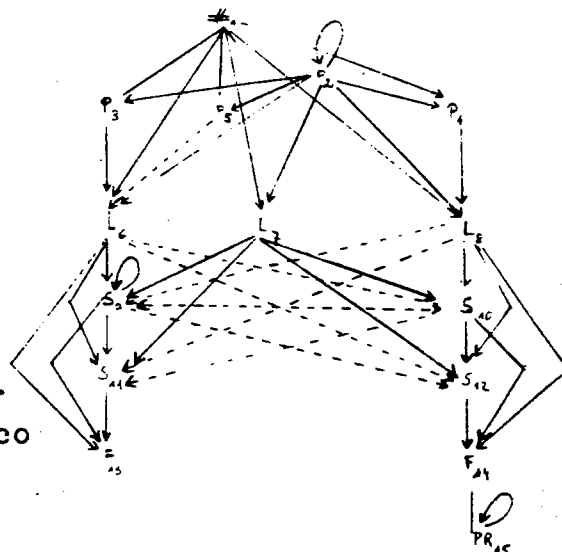
- 1 = inicio 2= prefijo recursivo 3=prefijo adnominal
 4= pref.adv. 5= prefijo denominativo 6= morfema no verbal
 7= morf.ligado 8= morfema verbal 9= suf.adnominal recursivo
 10=suf.adverbial recursivo 11= suf. " no-recursivo
 12=suf.adverbial no recursivo 13= flexion no verbal 14= pronombre

Esta red de transición esta expresada en reglas de producción del tipo:

premisa=complejo de posibles propiedades del morfema a tratar

condición=propiedades exigidas simultaneamente para activar la acción.

acción=transición a otro estado morfológico



Este saber operativo esta almacenado en bits -formato hexadecimal correspondientes a las propiedades morfológicas vistas arriba. Cada regla esta codificada como una estructura de datos separada e independiente. La regla siguiente permite la transición de morfema léxico no verbal L6 a su flexión correspondiente

i	0101 1110 0000 0000 0010 0001 1011 0000
ii	0100 0000 0000 0000 0000 0001 1011 0000
iii	0013 0002

De un estado x se pasa a F13 (flexión no verbal) si los bits 2,24,25,27,28 están prendidos. En caso afirmativo se comprueba que F13 y el morfema anterior sean del mismo subtipo (la condición 0002) ."ventan" exige "-a", y "as" es del tipo "a".Entonces se pasa al estado 13 o sea "as" recibe tal estado.

3 Componentes

Hay varios programas que:

1. preparaN y editan los datos : generan la codificación interna
: comprimen los datos (árboles)

A partir de estos formatos comunes para las 3 lenguas se pueden obtener listas comparativas de aspectos o categorías a partir de determinado vector. Estos serían productos secundarios a obtener de interes lingüístico.

2. Programas de entrada/salida según el proceso a adoptar a continuación. El formato del output será diferente según se quieran invertir o no los resultados lingüísticos, o se desee engranar una ATN.

3. Programas de procesamiento y control:

IBIS :generación y lectura de los registros para las tablas
ENTE :lectura de los datos de entrada
SEGMENT :segmentación del texto de entrada
KAKADU :comprobación de si la palabra es funcional o no
SPATZ :realización de transformaciones (lápices-->lápiz,...)
WACHTEL :Es el programa de descomposición morfológica y consta de las funciones: MORPH (busca arborea),DELTA (acepta o rechaza las transiciones)y RUECKSETZ

Hay varios posibles outputs de la descomposición morfológica, interno -a engranar con la lematización-, u otros según el proceso a continuar.

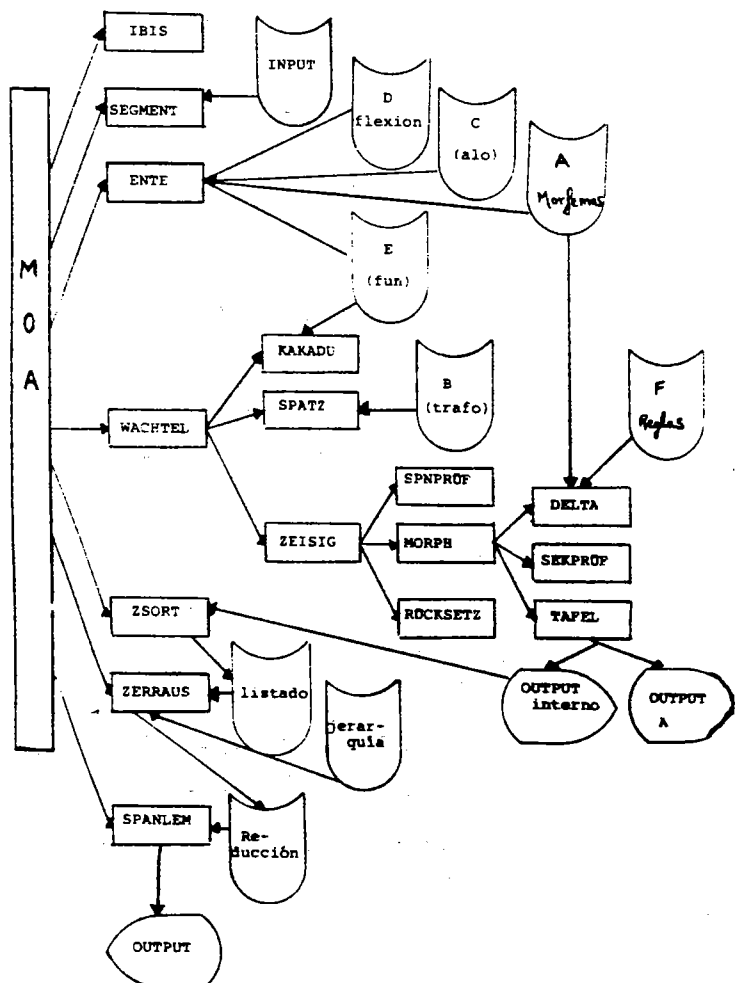
Ya que el mecanismo ofrece varias posibles descomposiciones para algunas palabras (por razones diacrónicas o fonéticas), éstas son alistadas según un rango de validez estadística (ZSORT), y luego lingüística (ZERRAUS) por la que se eligen las descomposiciones definitivas, si son varias en el caso de los homógrafos.

La lematización y categorización es un "matching" de los tipos y subtipos de los patrones morfológicos que tienen asignados cada morfema.

Diagrama del funcionamiento de los programas para la descomposición morfológica y lematización.

Este paquete de programas está implementado en SPL y ocupa unos 90 Kb (procesamiento y listas).

Actualmente existe una versión interactiva en un ordenador SIEMENS 7760.



"Prospecció Automatitzada de Textos Catalans"

El proyecto de investigación "Prospecció Automatitzada de Textos Catalans" (PATC), cuyas fases iniciales se están llevando a cabo en la Universidad de Barcelona, tiene sus orígenes en el deseo de crear un centro dedicado a la investigación lingüística asistida por ordenador. El Anteproyecto para la creación de un centro de lingüística automática (abril 1975) y la Memoria para la creación de un centro de lingüística automática en la Universidad de Barcelona (enero 1976), elaborados conjuntamente por el profesor Santiago Mollfulleda y por el que suscribe, son los primeros documentos redactados en relación con esta idea, que desde tiempo atrás estaba en la mente de sus autores.

Con motivo del primer intento de poner en marcha este proyecto, tuvimos que concretar mucho y delimitar nuestros objetivos. En relación con ello debemos citar por una parte la celebración en Barcelona el 26 y 27 de Mayo de 1977 de una mesa redonda sobre informática y lingüística a la que presentaron comunicaciones, aparte de los profesores E. Moreu, Juan Vernet, M^a Teresa Cabré y M^a Angels Vidal de Barcelona, prestigiosos investigadores de diversos centros europeos (L. Delatte, de Lieja, Mario Alinei, de Utrecht, M. Tournier, de St. Cloud, P. Imbs, de Nancy, G. Colon, de Basilea) —véase el nº 1 de la serie Informática y lingüística (Mesa redonda sobre informática y lingüística. Barcelona, 26 y 27 de Mayo de 1977), publicación ciclostilada editada por la "Fundación para el desarrollo de la función social de las comunicaciones" (FUNDESCO); y, por otra parte, debemos aludir a la memoria Prospecció Automatitzada de Textos Catalans (504 págs.) redactada según el esquema para memorias de trabajos científicos y proyectos de investigación de FUNDESCO cuando intentábamos llevar adelante el proyecto con el patrocinio de esta fundación y con el concurso de diversas instituciones catalanas (el equipo redactor de la memoria estuvo constituido por los profesores M^a Teresa Cabré, Carmen González, Teresa Gracia, Santiago Mollfulleda, Mercè Otero, Pere Quetgles y Joaquín Rafel, y por el informático Valentín Ibáñez). Esta po-

sibilidad no llegó a cuajar por diversos motivos, pero en aquel momento habíamos conseguido pasar de un proyecto inicial vago e impreciso a otro de un notable grado de concreción, descrito de un modo tan detallado que permitía el inicio inmediato de su ejecución.

Los medios de financiación no llegaron hasta el año 1980 en que se convocaron por primera vez las Ayudas a la investigación por la Universidad de Barcelona, en principio para proyectos de un año de duración. En los años sucesivos se han ido concediendo ayudas para diversas fases del proyecto.

Sin perder de vista el objetivo general que nos habíamos trazado al principio, el proyecto específico que ha resultado tiene por finalidad fundamental la de crear un archivo de textos —en principio literarios— de la lengua catalana sobre soporte magnético automatizable con la condición básica de mantener toda la información que se contiene en el documento llamado Información fuente (IF) (edición moderna, incunable, manuscrito, etc.), para su ulterior utilización en la investigación lingüística. Además, como objetivos accesorios, consideramos la creación de programas generales para la explotación científica de los textos almacenados. La gran mayoría de los textos están escritos en lengua catalana, pero algunos de ellos lo están en latín; se trata de textos medievales que tienen un gran interés para la historia de la lengua catalana.

Por lo que respecta a los textos en latín, se han introducido en soporte magnético y se han tratado informáticamente los Usatges de Barcelona y el Cançoner eròtic de Ripoll; para estos textos se han utilizado los servicios del Laboratoire pour l'analyse statistique des langues anciennes de la Universidad de Lieja.

En cuanto a los textos propiamente en lengua catalana, en estos momentos tenemos completamente introducidos y corregidos los textos siguientes:

Tirant lo Blanc (extensión: 3.000.000 de caracteres)
(IF: facsímil del incunable de 1492,
propiedad de la Hispanic Society of
America)

El Quadern Gris de Josep Pla (ext.: 1.600.000 car.)

Se está terminando de introducir la Obra narrativa completa de Joaquim Ruyra (ext.: 2.000.000 de caracteres) y se ha comenzado la introducción de las Memòries de Josep M^a de Sagarra (ext.: 2.500.000 caracteres). Para el año 1984 se tiene programada la introducción de la obra narrativa completa de Víctor Català (ext. 3.000.000 de car.) y de Narcís Oller (4.000.000 de car.).

La introducción de la información sobre soporte magnético automatizable se hace a través de una estación de datos IBM 3741, en disco flexible. El almacenamiento definitivo se hace en cinta magnética.

Aparte de la introducción en sí misma, una de las labores fundamentales del equipo de trabajo ha sido la elaboración de los distintos documentos relacionados con el proceso informático, empezando por las codificaciones e instrucciones de entrada de datos y de corrección, y terminando por los programas para la obtención de distintos índices, aparte de otros trabajos que se están realizando.

Los índices que pensamos obtener de todo texto introducido —ya existen los de El Quadern Gris— son:

1º Índice de palabras ordenadas alfabéticamente con indicación de las frecuencias absoluta y relativa de cada una.

2º Índice de palabras ordenadas por orden decreciente de frecuencias.

3º Índice de palabras ordenadas alfabéticamente, acompañadas de la indicación de las referencias de su ubicación en el texto (página, línea, número de orden de la palabra en la línea —eventualmente, columna, capítulo, etc.).

4º Índice inverso (palabras ordenadas por la parte final), con indicación del número de finales distintos de uno, dos y tres caracteres).

Frecuentemente, en vez —o además— del nº 3, se obtendrá lo que se conoce con el nombre de concordancias, es decir, además de la información que contiene el índice nº 3, y junto a la referencia de cada ocurrencia, un fragmento del texto para cada

una de las apariciones de cada palabra.

Para todos estos índices hemos elaborado los programas adecuados y los resultados son listados sobre papel y almacenados en cinta magnética (una muestra de ellos, por lo que respecta a la obra El Quadern Gris, son adjuntados a este informe). Existe el proyecto de hacer una edición de ellos en microficha para utilizar a través de una lectora-copiadora.

Al margen de estos resultados particulares, o de otros cualesquiera, el texto queda almacenado, tanto en su primitiva forma, que reproduce la de la Información fuente, como en forma de lo que hemos llamado base de datos. Esta base de datos consta de un registro de 26 campos para cada palabra de la obra; en él se encuentra, aparte de la misma palabra, toda la información necesaria para su adecuada referenciación, catalogación, ordenaciones diversas, etc. A la base de datos se puede acudir para obtener cualquier información que no esté contenida en los índices editados.

Además del desarrollo de los aspectos mencionados del proyecto, constituyen también parte del mismo los trabajos que se están llevando a cabo para realizar una lematización automatizada de los textos. La asociación de todas las formas flexivas a una única forma representativa de la palabra (LEMA), y la separación de formas homógrafas, asociándolas a lemas distintos, es uno de los caballos de batalla del tratamiento exhaustivo de textos. Nuestro proyecto de lematización comprende un programa —todavía no completamente elaborado— que obtendrá un lema (o varios posibles lemas) después de someter cada palabra del texto a una serie de substituciones de su parte final, alternadas con la confrontación del resultado de cada operación con un diccionario almacenado sobre soporte magnético automatizable; los resultados de todas estas operaciones serán sometidos a un investigador para que los dé por buenos o para que elija entre varios lemas propuestos. En el momento actual tenemos sobre soporte magnético un diccionario de 50.000 entradas con sus indicadores gramaticales normalizados y elaborados los algoritmos de análisis morfológico para el tratamiento de plurales de sustantivos,

femeninos y plurales de adjetivos, flexión verbal, diminutivos y superlativos sintéticos.

Por lo que se refiere a los aspectos materiales (hardware), —aparte de la utilización del centro de Lieja para los textos en latín, a que ya nos hemos referido— unas fases del proyecto se han desarrollado en un centro de cálculo comercial, cuando las condiciones del Laboratorio de Cálculo de la Universidad no reunían los requisitos mínimos exigidos para un trabajo de este tipo; desde hace un año, tras la instalación de una máquina IBM 4341 y una reorganización del Laboratorio de Cálculo, con instalación de terminales en las Facultades, etc., hemos empezado a utilizar estos servicios. El principal problema que tenemos pendiente es el de la obtención de una impresora adecuada para los resultados en papel; de momento, tal como se ve en los especímenes que reproducimos junto a este documento, hemos debido recurrir a una ficción para representar los conceptos más imprescindibles, con impresión en interlínea (acento grave, acento agudo, diéresis, mayúscula), o con sobreimpresión (cedilla). En los listados de trabajo (correcciones, etc.) estos y otros conceptos vienen representados por su referencia interna y se trabaja con listas de correspondencias, pero en los resultados para ofrecer a personas externas al propio proyecto, nos pareció que, aunque fuera provisionalmente, debíamos elaborar los programas que permitieran una lectura no mediatizada por instrucciones demasiado complejas. De todas formas, no olvidamos el problema, y estamos colaborando con la dirección del Laboratorio de Cálculo en el sentido de estimular la instalación de una impresora adecuada que permita representar cualquier tipo de signo especial.

Ni que decir tiene que el futuro de un programa de esta naturaleza depende en gran medida de las posibilidades de financiación, y que el modo como hasta ahora se ha venido desarrollando este proyecto no asegura su continuidad. Esperamos —y trabajamos en ello— poder consolidar este aspecto fundamental y poder continuar ofreciendo información sobre futuros logros.

Joaquim RAFEL i FONTANALS

1. Índice de palabras ordenadas alfabéticamente con indicación de las frecuencias absoluta y relativa de cada una.

JARDEN 3115								274
MOT		F.ABS	F.REL	MOT		F.ABS	F.REL	
POSAVA		24	0,009	POSSEIR		8	0,003	
POSAVEN		3	0,001	POSSEIT		1	0,000	
POSAVEN		5	0,001	POSSESSIC		9	0,003	
POSEM		6	0,002	POSSESSIVES		1	0,000	
POSEN		13	0,005	POSSIBILISME		1	0,000	
POSES		4	0,001	POSSIBILITAT		23	0,009	
POSSESSIN		1	0,000	POSSIBILITATS		15	0,005	
POSEU		3	0,001	POSSIBLE		101	0,039	
POSI		5	0,001	POSSIBLEMENT		1	0,000	
POSICIO		43	0,016	POSSIBLES		8	0,003	
POSICIONS		6	0,002	POST		2	0,000	
PCSIN		1	0,000	POSTA		9	0,003	
POSIS		1	0,000	POSTAL		3	0,001	
PCSITIU		9	0,003	POSTAL		1	0,000	
POSITIUS		2	0,000	POSTALS		1	0,000	
POSITIVA		16	0,006	POSTERIOR		7	0,002	
POSITIVAMENT		11	0,004	POSTERIORITAT		2	0,000	
POSITIVES		6	0,002	POSTERIORES		4	0,001	
POSITIVISME		1	0,000	POSTES		2	0,000	
PCSG		13	0,005	POSTISSA		1	0,000	
POSSEIIX		9	0,003	POSTRES		7	0,002	
POSSEIIXEN		1	0,000	POSTULEN		1	0,000	
POSSEIIXI		1	0,000	POSTURES		1	0,000	
POSSEI		1	0,000	POSTURETES		1	0,000	
POSSEIA		1	0,000	POT		225	0,088	
POSSEIDES		1	0,000	POTA		1	0,000	
POSSEIDOR		1	0,000	POTABLE		1	0,000	
POSSEIDORES		1	0,000	POTABLES		2	0,000	
POSSEIEN		2	0,000	PCTACURT		1	0,000	

2. Índice de palabras ordenadas por orden decreciente de frecuencias.

----- 39		
MOT	F.ABS	F.REL
ESPERANÇA	11	0,004
ESSERS	11	0,004
ESTAVEN	11	0,004
ESTRIS	11	0,004
ESTUDIS	11	0,004
EXACTA	10	0,003
EXAMEN	10	0,003
EXAMINAR	10	0,003
EXCEPCIONAL	10	0,003
EXERCICIS	10	0,003
EXPLICACIÓ	10	0,003
EXPOSICIÓ	10	0,003
EXQUISIDA	10	0,003
EXQUISIDES	10	0,003
EXQUISIT	10	0,003
FANAL	10	0,003
FATXENDA	10	0,003
FEBRE	10	0,003
RASTRE	10	0,003
FINAL	10	0,002
FÍSICAMENT	10	0,003
ABUNDANCIA	10	0,003
FORASTERS	10	0,003
ACLAPARAT	10	0,003
RELACIONS	10	0,003
ACTIVA	10	0,003

3. Índice de palabras ordenadas alfabéticamente acompañadas de las referencias de su ubicación en el texto (página, línea y número de la palabra dentro de la línea en la Información fuente).

QUADERN GRIS	-----									
401-----	LOCALITZACIONS-----									
DESTACA VA	579	4	3,	850	17	10,				
DESTAPEN	299	32	6,							
DESTAQUEN	664	26	5,							
DESTEIXIR	406	17	10,							
DESTEMPRADA	333	5	7,							
DESTENYIDES	503	12	4,							
DESTENYIT	542	27	2,							
DESTI	295	12	7,	692	36	6,				
DESTIL·LACIÓ	404	17	6,							
DESTIL·LAT	511	12	4,							
DESTIL·LO	429	17	5,							
DESTINADA	316	5	4,	326	18	1,	365	22	5,	455
	808	33	6,	843	34	5,	849	12	1,	
DESTINADES	274	29	6,	535	5	6,	854	33	6,	
DESTINAREN	431	15	1,							
DESTINAT	324	11	4,	390	32	4,	474	18	7,	741
										30
										8,
										838
										23
										1,
DESTINATS	356	3	7,	569	39	6,	814	19	1,	
DESTORBADA	182	30	5,							
DESTRALEJAR	394	37	2,							
DESTRALS	399	18	8,							
DESTREMPADA	550	23	9,							
DESTREMPADES	503	34	3,							
DESTRIAR	761	13	3,							
DESTRUSSA	594	20	6,							
DESTROSSAT	572	27	8,							
DESTRUCCIÓ	220	39	9,	295	21	7,	303	35	2,	456
	483	32	2,	526	28	8,	817	11	3,	819
										38
										7,
DESTRUEIX	746	10	2,							
DESTRUÏDA	536	8	4,							

4. Índice inverso (palabras ordenadas alfabéticamente por su parte final, con indicación numérica de los finales distintos de los tres, dos y un caracteres).

CUADERNO GRIS
FA---XXX---FA---XX-----MOT-INV---FA---XXX---FA---XX-----MOT-INV----- 24

		SORRUDA				GALLOFA
		MOL SUDA				SOFA
-UDA 38	-DA 877	PANACEA	-QFA 4			DESFA
		OCEA				SATISFA
-CEA 2		IDEA	-SFA 2			BUFA
-DEA 1		FARMACPEA				BUFA
		EUROPEA				BALDUF
-PEA 2		AREA				ESTUFA
		CREA	-UFA 4	-FA 24		EMBRIAGA
		MENYSPREA				BALIGA-BALAGA
		VERBORREA				ARGELAGA
-REA 4		ODISSEA				PLAGA
		NAUSEA				AMAGA
-SEA 2	-EA 11	FA				PASTANAGA
-FA 1		EMBAFA				PAGA
		AGAF				APAGA
		AGAF				NAUFRAGA
		GIRAF				NISSAGA
-AFA 4		REFA				VAGA
-EFA 1		RIFA				DIVAGA
		TARIFA				MANYAGA
		TIFA	-AGA 13			CEGA
		REVIFA				ENCEGA
-IFA 4		PELFA				OFEGA
		SOLFA				MARFEGA
-LFA 2		LIMFA				SACRILEGA
		NIMFA				COLLEGA
-MFA 2		COFA				PLEGA
		FOFA				REPLEGA



ETIQUETADO NO FONEMATICO PARA UN SISTEMA AUTOMATICO MONOLOCUTOR DE =====

RECONOCIMIENTO DE PALABRAS AISLADAS =====

Por Antonio J. Rubio Ayuso (*) y M. Carmen Carrion Perez (**)

(*) Departamento de Electronica

(**) Departamento de Electricidad

Universidad de Granada

I. Introduccion -----

Presentamos en estas lineas una de las partes de nuestro trabajo que pueden considerarse terminadas desde el punto de vista logico. Trabajamos en otros temas dentro del mismo campo, pero eso lo presentaremos en otros numeros.

Aqui presentamos un Sistema Automatico de Reconocimiento Monolocator con Etiquetado No Fonemático (SARMENF) (1). Funciona a base de dividir la señal sonora en segmentos de la misma duracion e intentando clasificar cada segmento entre unos prototipos obtenidos durante la Fase de Entrenamiento, mediante un criterio de distancia euclídea. La etiqueta que se asigna a cada segmento no guarda ninguna relacion con los fonemas del lenguaje. Este tipo de etiquetado puede ser ventajoso debido a la difícil localización del fonema como unidad de analisis.

II. Esquema general -----

Hay dos fases en el funcionamiento del SARMENF: fase de entrenamiento y fase de reconocimiento.

2.1 Fase de Reconocimiento -----

Durante la Fase de Reconocimiento se realizan los siguientes pasos:

1) Una vez que el locutor ha pronunciado una palabra del vocabulario permitido, el primer paso consiste en muestrear la señal, in-



introduciendo las muestras en un ordenador.

La frecuencia de muestreo es de 10 kHz. La conversión A/D se hace a través de un conversor A/D de 12 bits. Para cada palabra, se toman 15000 muestras, durante 1'5 s.

2) Con objeto de reducir el volumen de cálculo, los extremos de la señal sonora son localizados. En general, ninguna palabra ocupa completamente las 15000 muestras.

3) Una vez que los extremos han sido localizados, se establece una segmentación regular de la señal sonora. Los segmentos son de 10 ms (100 muestras) y se solapan a la mitad.

Sobre cada segmento se calcula un conjunto de parámetros. En primer lugar se hace una clasificación previa entre segmentos sonoros y no sonoros, utilizando la Energía y el Número de pasos por cero del segmento (ZCR). Se clasifican como segmentos sonoros aquellos que cumplen la inequación

$$\text{Energía} > A+B \cdot \text{ZCR}$$

donde las constantes $A = .1605E+08$ y $B = 1.075E+06$ se han determinado experimentalmente y son función de la ganancia global del sistema de audio.

Se clasifican como no sonoros aquellos segmentos que cumplen la desigualdad

$$\text{Energía} > C+D \cdot \text{ZCR}$$

donde las constantes $C = .2012E+08$ y $D = 8.047E+05$ también se han determinado experimentalmente.

Aquellos segmentos que no cumplen ninguna de las dos condiciones anteriores son clasificados como dudosos y son ignorados en el tratamiento posterior.

Si un segmento se clasifica como sonoro, se calculan sobre el otros cuatro parámetros adicionales: los tres primeros formantes y el logaritmo de la razón de las áreas de la primera y segunda sección del tubo acústico equivalente al tracto vocal. Esta determinación se hace mediante Predicción Lineal (2).

Por otra parte, si el segmento se clasifica como no sonoro, los cuatro parámetros adicionales son los valores medios del espectro de la señal dividido en cuatro zonas o bandas (0 - 1.25 kHz, 1.25 - 2.5 kHz, 2.5 - 3.75 kHz y 3.75 - 5 kHz).

4) Los parámetros correspondientes a un segmento se comparan con los de los Segmentos de Referencia obtenidos en la Fase de Entrenamiento. De esta comparación se obtiene una etiqueta para el segmento en cuestión. Así se consigue una cadena de etiquetas (una etiqueta por cada 5 ms) que representa a la palabra emitida.

5) El último paso en la Fase de Reconocimiento consiste en comparar la cadena de etiquetas con las Cadenas de Referencia, obtenidas también en la Fase de Entrenamiento (una Cadena de Referencia por cada palabra del vocabulario). Esta comparación se lleva a cabo por medio de un algoritmo de Programación Dinámica (3). Este algoritmo permite comparar cadenas de etiquetas de diferente longitud, pro-



porcionando la "Similitud" entre ellas en funcion de la distancia euclidea entre los diferentes simbolos de cada una de las cadenas comparadas.

Un esquema de bloques para la Fase de Reconocimiento puede verse en la parte derecha de la Figura 1.

2.1 Fase de Entrenamiento

Esta Fase esta representada en la parte izquierda de la Figura 1.

El Entrenamiento consiste en la pronunciacion de todas las palabras del vocabulario, cada una de ellas repetida cuatro veces.

Despues del muestreo de cada palabra de entrenamiento, se calculan sus extremos y se obtienen los parametros correspondientes a los segmentos de la señal, del mismo modo que en la Fase de Reconocimiento.

Cada segmento se considera como un punto de un hiper-espacio euclidiano de seis dimensiones. En realidad se manejan dos hiper-espacios: uno para los segmentos sonoros y otro para las no sonoros. De este modo, la voz del locutor se caracteriza por una cierta distribucion de puntos a lo largo de ambos hiper-espacios.

Los pasos que siguen en la Fase de Entrenamiento son:

1) Analisis de las nubes formadas por los segmentos de las palabras pronunciadas, en ambos hiper-espacios.

Aquellas nubes claramente separadas del resto de puntos se consideran como correspondientes a segmentos de señal del mismo tipo. Su centro geometrico representara al Segmento de Referencia, con una etiqueta asociada.

La obtencion de las nubes y Segmentos de Referencia se lleva a cabo mediante el algoritmo "k-medias" modificado (4).

La modificacion puede describirse brevemente. La principal fuente de fallos en el algoritmo "k-medias" original esta en la seleccion del conjunto inicial de centros. Evitamos este problema de la siguiente forma:

Se hace una normalizacion de todos los parametros de modo que la media de cada parametro sea 1. Asi, el punto (1,1,1,1,1,1) es el centro de la distribucion completa. Entonces, se colocan los puntos iniciales de modo que esten uniformemente repartidos alrededor del punto (1,1,1,1,1,1).

El numero de Segmentos de Referencia (k) se ha establecido experimentalmente. Para los segmentos sonoros $k = 21$ y para los segmentos no sonoros $k = 7$.

2) De acuerdo con los Segmentos de Referencia, cada palabra del vocabulario se analiza con el fin de elaborar una cadena de etiquetas (Cadena de Referencia) para esa palabra. Estas Cadenas de Referencia son las que se usan en el algoritmo de Programacion Dinamica (Fase



de Reconocimiento) para comparar una palabra emitida con todas las del vocabulario.

III. Conclusion

En este trabajo se ha presentado un sistema automatico de reconocimiento monolocutor con etiquetado no fonemático. Las principales características son:

1) Es un sistema de reconocimiento monolocutor para palabras aisladas.

2) Se lleva a cabo un etiquetado no fonemático. Hay 21 diferentes tipos de segmentos sonoros y 7 tipos de segmentos no sonoros.

3) El vocabulario es un conjunto de 19 palabras castellanas relacionadas con el calculo aritmetico: "Cero", "Uno", "Dos", "Tres", "Cuatro", "Cinco", "Seis", "Siete", "Ocho", "Nueve", "Parentesis", "Cierra", "Mas", "Menos", "Por", "Partido", "Raiz", "Igual?" y "Punto". Este vocabulario permitira el manejo oral del ordenador en calculos aritmeticos simples.

4) Los resultados experimentales estan hasta el momento alrededor del 65% de reconocimiento correcto. Merece la pena destacar que no se ha puesto especial atencion en la seleccion de las Cadenas de Referencia.

5) El sistema ha sido implementado sobre un ordenador ECLIPSE-250 de DATA GENERAL, en lenguaje de programacion FORTRAN V.

IV. Referencias

- (1) R. Reddy
"Speech recognition by machine: a review"
Proc. IEEE, vol. 64, pp. 505-531
 - (2) J. D. Markel, A. H. Gray
"Linear Prediction of Speech"
Springer-Verlag, 1976.
 - (3) H. Sakoe, S. Chiba
"Dynamic programming algorithm optimization for spoken word recognition"
IEEE Trans. on ASSP. Vol. ASSP-26, n. 1, pp. 43-49, 1978.
 - (4) J. T. Tou, R. C. Gonzalez
"Pattern recognition principles"
Addison-Wesley, 1974.
-

Nota: Este trabajo ha sido presentado al "1983 SPAIN WORKSHOP ON SIGNAL PROCESSING AND ITS APPLICATIONS" que se celebrara en Sitges (Barcelona) en Septiembre de 1983.

ESQUEMA - RESUMEN

Martín S. Ruipérez, Catedrático de Filología Griega
Universidad Complutense, Madrid

Aparte de temas de Filología Clásica y Lingüística Comparada Indoeuropea ha trabajado en:

- Lingüística estructural y funcional (griego y español), concretamente morfosintaxis y fonología.
- Con equipo de alumnos y colaboradores, sin ordenador y sólo con calculadoras, ha establecido (1968) un vocabulario de frecuencias de griego clásico de 400-350 a. C. Igualmente un índice morfológico verbal de frecuencias. Todo ello para la elaboración de un método de enseñanza del griego antiguo, que está en fase de experimentación en varios Institutos y Facultades.
- Ha dirigido una tesis doctoral de lingüística cuantitativa aplicada al problema de los arcaísmos en Homero.
- Desde 1973 se ha familiarizado con cálculo programado.
- Desde 1980 se ha familiarizado con ordenadores, concretamente con un HEWLETT-PACKARD HP250, provisto de base de datos.

Además de programas comerciales, ha realizado programas lingüísticos:

1. Construcción de cadenas a partir de segmentos conmutables. Concretamente, conversión de un número introducido en dígitos en texto español, con perfecta coherencia sintáctica.

2. Corte de palabras conforme a las reglas de la Academia y al buen uso tipográfico para programas de textos cuando al final de una línea no cabe la palabra entera. Versiones para Castellano, catalán, francés e italiano. Disponible en el mercado.

Tiene un programa de calendario (que incluye la reforma del calendario gregoriano) para determinar en qué día de la semana cae cualquier fecha entre el 1 de enero del año 1 (!) y el 28 de febrero del 2100; número de días entre dos fechas y fecha a partir de otra y de un número de días de intervalo.

Todo en lenguaje BASIC y "Extended BASIC" o "Business BASIC" (en realidad BASIC enriquecido con PASCAL)

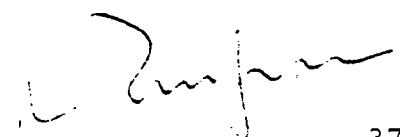
- Desde diciembre de 1982 dispone en su cátedra de un modesto equipo compuesto por:

- Ordenador HEWLETT-PACKARD HP87 con 128 Kb de memoria RAM, más un ROM de Input/Output y otro de Assembler.
- Periférico de una unidad dual de diskettes floppies con capacidad total de 540 Kb
- Periférico de impresora de agujas EPSON III/100

Actualmente con este equipo está desarrollando "software" para tratamiento científico de textos: introducción, almacenamiento, modificación, impresión, búsqueda y confección de índices alfabéticos (con alfabetización distinta para el griego, en el que, por ejemplo, la z no va al final sino después de la e).

Se prevé la aplicación inmediata al análisis de textos métricos griegos codificados y de un repertorio de inscripciones.

- En 1982-83 curso impartió la asignatura de Lingüística Griega en la Universidad de Sevilla, 15 de abril de 1983. *Lingüística griega y filología*



USL: UN SISTEMA PARA INTERROGAR EN CASTELLANO

A BASES DE DATOS RELACIONALES

Luis de Sopena.

Es un sistema experimental desarrollado en el Centro Científico de IBM en Heidelberg (Alemania) inicialmente por H. Lehmann, N. Ott y M. Zoeppritz. Información básica sobre él puede encontrarse en (3, 4). Se han implementado tres versiones de USL: alemana, inglesa y española, que es la que aquí vamos a considerar (5).

1. OBJETIVOS

USL constituye una interfase interactiva con un Sistema de Gestión de Bases de Datos Relacionales. Traduce las frases de entrada escritas en Lenguaje Natural a sentencias del lenguaje formal de interrogación de la base de datos. Con dicha expresión accede a la Base, y una vez obtenida la respuesta, se la presenta al usuario, bien directamente, bien después de llevar a cabo alguna operación adicional sobre ella.

El principio fundamental de diseño de USL es la independencia respecto del dominio de la aplicación, es decir que el sistema es utilizable para diferentes dominios. Para lograr esto han debido desarrollarse las reglas de la gramática española y las rutinas de interpretación adecuadas para el subconjunto del LN que se necesita para manipular las informaciones de una base de Datos.

USL se dirige a usuarios no especialistas en Informática, pero con un cierto conocimiento del dominio de la aplica-

ción:..

Por otra parte, no intenta simular diálogos ni construir modelos de adquisición de conocimientos.

2. DESCRIPCION DEL SISTEMA

El método utilizado para el proceso de traducción mencionado se basa en el análisis sintáctico y semántico del subconjunto del castellano escogido.

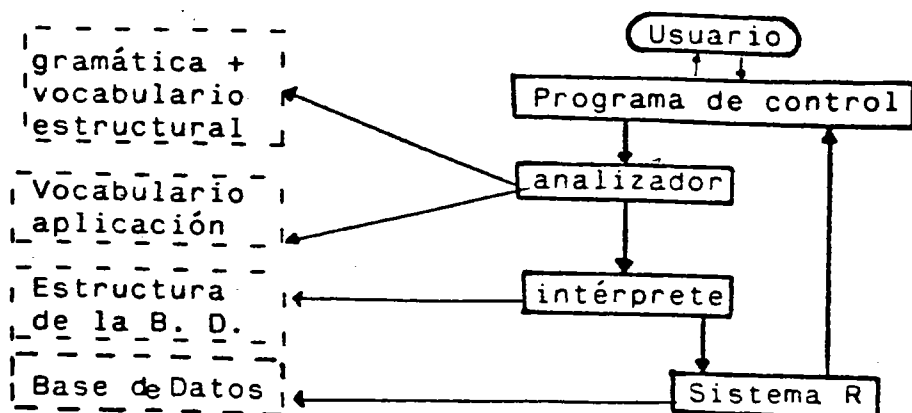
El sistema se compone de:

- Un analizador llamado USAGE (2), que acepta gramáticas escritas en un formato BNF modificado, y permite la utilización de rutinas semánticas.
- Una gramática de unas 450 reglas, donde se codifican las estructuras sintácticas permitidas.
- Un vocabulario de unas 400 palabras estructurales, incluyendo cuantificadores, preposiciones, pronombres, nombres de meses, los verbos tener y ser, etc.
- Un conjunto de rutinas de interpretación, especificadas en las reglas de la gramática, que llevarán a cabo de traducción al lenguaje de la base de datos.
- Un sistema de Base de Datos Relacionales: el sistema R (1), que dispone de su propio lenguaje formal de interrogación, SQL.

Además, el usuario deberá definir para cada aplicación particular:

- Las relaciones base que constituyen la base de datos, y las relaciones virtuales (visualizaciones) definidas sobre ellas.
- El vocabulario asociado con la aplicación.

El diagrama del sistema es:



Las palabras estructurales y el significado de las estructuras sintácticas se suponen independientes del dominio y son en consecuencia interpretadas por las rutinas semánticas de un modo general. Por el contrario, los nombres, verbos y adjetivos se consideran dependientes de la aplicación, y deben de ser definidos por el usuario de acuerdo con la estructura de las relaciones de la Base de Datos. Los nombres de contenidos y los valores numéricos no necesitan ser definidos.

Los criterios seguidos para la representación del conocimiento en USL son los siguientes: el conocimiento conceptual se encuentra contenido en las relaciones virtuales, que están definidas sobre las relaciones físicamente almacenadas en la base, y ligadas a las palabras del vocabulario; el conocimiento factual está almacenado directamente en la base de datos, y no ha sido duplicado en forma de diccionario o de red semántica.

En USL hay una separación total entre el analizador y la base de datos: durante el análisis no se accede a la base, durante la interpretación solamente se utiliza información sobre su estructura.

3. COBERTURA DEL LENGUAJE

Para describir la cobertura del castellano que el sistema proporciona usaremos ejemplos clasificados según la estructura o construcción sintáctica:

- Interrogativas generales: Pertenece Holanda a la OTAN?
- Interrogativas parciales: Qué país exporta petróleo?
- Ordenes : Lista los países árabes
- Enunciativas: Juan vive en Madrid
- Complementos del nombre : Qué empleados de Juan viven en Madrid?
- Negación : Qué países no importan petróleo?
- Adjetivos: Cuales son los empleados solteros.
- Aposiciones: Qué empleados de García, el jefe, están casados?
- Oraciones de relativo: Cuáles son los países que exportan vino?
- Coordinación: Qué edad y qué salario tienen los empleados; quién vive en Madrid o en Barcelona?
- Cuantificadores: Tienen todos los jefes al menos 5 empleados?
- Comparación : Quien gana más que Juan?
- Posesivos: Qué empleado gana más que su jefe.
- Complementos de tiempo: Quién nació el 5 de marzo de 1.959?
- Funciones de suma, media, máximo, mínimo: Cuál es el salario medio?
- Subtotales: Cuál es la suma de los salarios de los empleados de cada jefe?

4. APLICACIONES Y RESULTADOS

En su estado actual, USL puede ser directamente utilizado para interrogar y actualizar una Base de Datos Relacional, y los usuarios pueden formular un número virtualmente ilimitado de preguntas diferentes utilizando el castellano como lenguaje de comunicación.

Aunque la versión española de USL aún no ha sido sometida a tests utilizando usuarios y aplicaciones reales, se han construido dos bases de datos experimentales para la puesta a punto y demostración del sistema, una con datos geográficos y otra con datos sobre los empleados de una empresa. Los ejemplos antes mencionados se han tomado de dichas bases.

Son interesantes, la independencia del sistema respecto del dominio, el hecho de que sea una interfase con un sistema general de Bases de Datos de estructura relacional, así como las facilidades de uso y definición de palabras.

REFERENCIAS

- /1/ ASTRAHAN M. M. et. al.
"System R: Relational approach to database management"
ACM Transactions on Database Systems, vol. 1, nº 2,
junio 1976.
- /2/ BERTRAND O., J. DAUDENARDE
"Usage-processing and natural language"
IBM Paris Scientific Center, 1977.
- /3/ LEHMANN H.
"Interpretation of natural language in an information
system", IBM Journal of Research and Development,
vol. 22, 1978.
- /4/ OTT N., ZOEPFRITZ M.
"USL - An experimental information system based on
natural language", in Bolc (Ed.): Natural Communication
with computers, vol. 2, Carl Hanser Verlag, Muenchen-
Wien, 1979.
- /5/ SOPEÑA L.
"Grammar of Spanish for the user specialty languages
system", IBM Heidelberg Scientific Center, 1982.

Q-SENECA (Un sistema pregunta respuesta en castellano)

M. F. Verdejo - Facultad de Informática, Universidad del País Vasco

Presentación

Desde la perspectiva de la inteligencia artificial, el estudio del lenguaje natural tiene dos objetivos:

1) Facilitar la comunicación con el ordenador para que accedan al mismo usuarios no especialistas

2) Modelar los procesos cognoscitivos que entran en juego en la comprensión del lenguaje para diseñar sistemas que realicen tareas lingüísticas complejas (traducción, resúmenes de textos, etc.)

Hay problemas en los que interesa fundamentalmente el primer objetivo. Lo que se desea es conseguir un intérprete para una clase de aplicaciones en un dominio restringido, que haga de traductor entre el ordenador y el usuario. El intérprete realiza dos tareas: una, de reconocimiento de la pregunta del usuario, otra de generación de una expresión equivalente en el lenguaje formal que utilice el ordenador para la aplicación. Este enfoque modela el lenguaje como una herramienta de comunicación (usuario-sistema) sobre conjuntos de información tipificada y restringida y, por tanto, sólo tratará el subconjunto del lenguaje natural que describirá los aspectos significativos en ese dominio. La comprensión de una consulta se plantea de este modo como un proceso de "detección" de un texto, de la información relevante.

El segundo objetivo se plantea el lenguaje como objeto de estudio, y la comprensión como un proceso complejo en el que intervienen grandes cantidades de conocimiento de naturaleza diferente (morfología, sintaxis, semántica, pragmática) y mecanismos de tratamiento variados (de comparación, búsqueda, inferencia aproximada, deducción...).

Algunos de estos aspectos son también objeto de estudio por parte de lingüistas, psicólogos, lógicos y filósofos, y por ello las aportaciones de unos y otros son fundamentales en la elaboración de teorías que den paso a la formalización de algunos fenómenos del proceso de comprensión. Desde la Inteligencia Artificial se aborda el problema con una óptica propia: la teoría debe permitir construir un sistema automático, que no tiene por qué estar basado necesariamente en una simulación del procesador humano.

El propósito de nuestro trabajo es la construcción de un sistema pregunta-respuesta que acepta consultas en castellano y es capaz de responderlas.

Un sistema pregunta-respuesta se compone de dos fases. La primera es un proceso de análisis o comprensión de la pregunta. La segunda, es un proceso de resolución o construcción de la respuesta.

La potencia de un sistema pregunta-respuesta depende fundamentalmente del modelo de representación del conocimiento sobre el que está basado. Nuestro sistema se basa en un modelo general en forma de red semántica: SENECA (G. Camarero & G. Sanz & M. F. Verdejo 79 a,b). Para cada aplicación permite representar el conocimiento específico, dotándole de una estructura que facilita la búsqueda y deducción de información, procesos esenciales en la comprensión de la pregunta y generación de la respuesta.

Este método de representación es aplicable a diferentes dominios y en particular lo hemos utilizado para construir un sistema automático de consulta a un banco de datos formado por textos literarios, de los que se quiere resultados estadísticos, usualmente requeridos en lingüística computacional.

En este dominio, el conocimiento se descompone en las siguientes unidades conceptuales:

Objetos: obras literarias, y en general, textos que representamos por nodos



Propiedades: de una obra literaria: tener título, autor, ficha, editorial, etc. De un texto: su codificación, índice... Las representaremos por nodos



Acciones: operaciones que pueden efectuarse sobre los objetos, por ejemplo, hallar frecuencias, contextos ... y las representamos por nodos



Estas unidades están relacionadas entre sí, formando la red de la figura 1, que constituye el conocimiento conceptual que el sistema tiene acerca del dominio.

Caracterización

El subconjunto del castellano que tratamos comprende las frases que se utilizan normalmente en consultas a bases de datos. Consideramos tres tipos. El primero, está formado por las consultas acerca de los objetos del dominio, ya sea para preguntar sobre su existencia, sobre el número de objetos que verifican una propiedad o relación o bien la especificación de un objeto del que se da una descripción parcial. Ejemplos son las preguntas:

Pregunta 1. ¿Hay alguna obra literaria del Siglo de Oro?

Pregunta 2. ¿Cuáles son los adjetivos que aparecen en las Soledades?

Pregunta 3. ¿Hay alguna palabra de frecuencia mayor de 100 en la Realidad y el Deseo?

Pregunta 4. ¿Cómo está codificado el libro del Buen Amor?

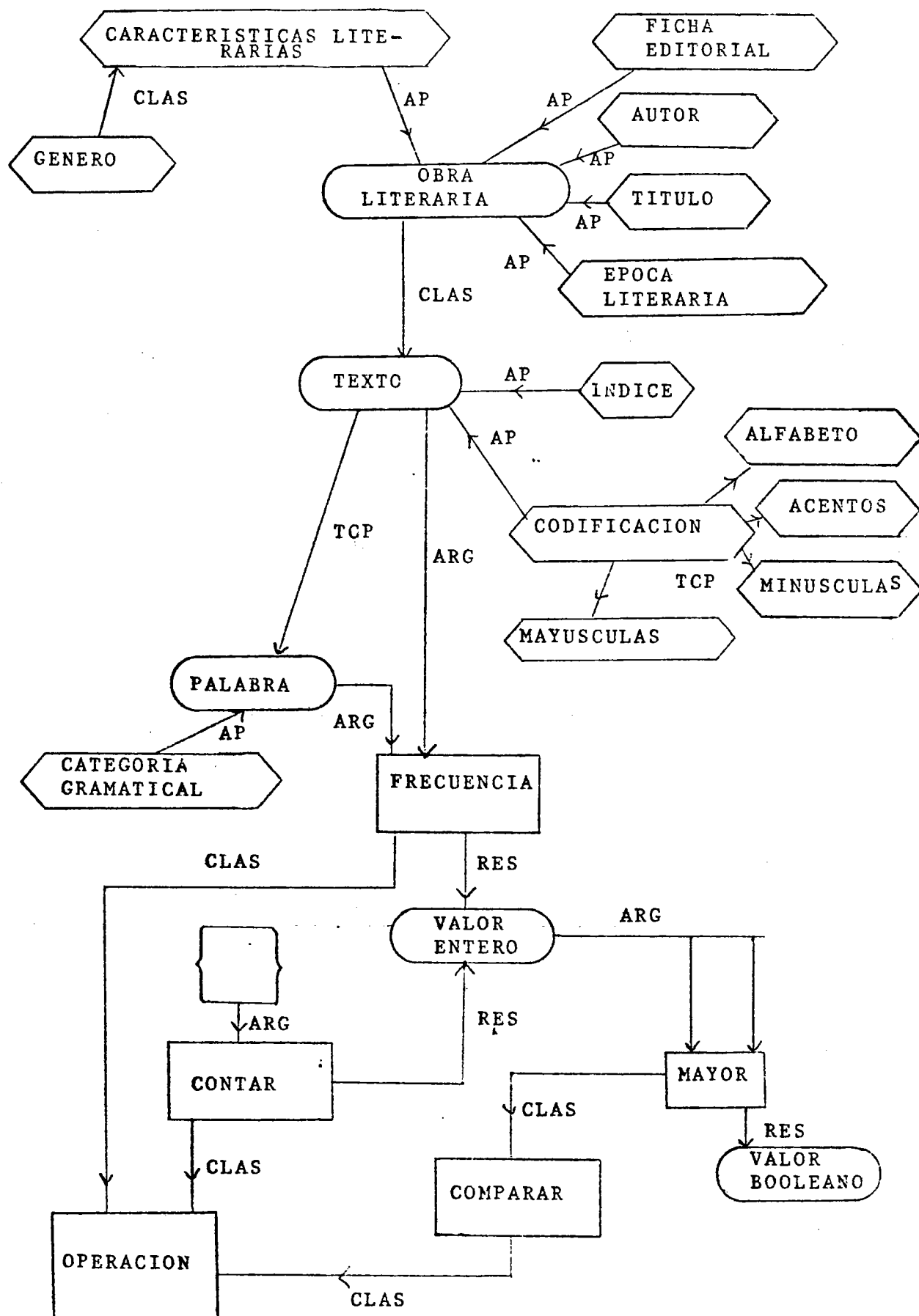


Figura 1.

Análisis de la consulta

El analizador acepta una consulta cuando es capaz de representarla en forma de red semántica. En este proceso entran en juego diferentes tipos de conocimiento: sobre el léxico, reglas sintácticas y semántica del lenguaje, por una parte. Sobre la propia representación por otra. Nuestro analizador tiene como modelo semántico el plano conceptual de la red del dominio problema. Una parte del léxico está asociado mediante un diccionario a elementos (nodos o subredes) de dicha red.

El objetivo del analizador es extraer el significado de la consulta, es decir descomponerla en términos de la red del dominio y representarla en forma de red semántica.

En la consecución de esta tarea el papel que juega el conocimiento que se tiene sobre el dominio es fundamental, no solo para verificar que la pregunta que se formula es coherente con él mismo, sino también para levantar posibles ambigüedades. Por esta razón el analizador interacciona continuamente con la red dominio, su modelo semántico, que le permite dirigir y orientar el análisis.

El proceso se basa sobre criterios semánticos, en el sentido de que se intenta detectar el elemento conceptual sobre el que se pregunta, y una vez localizado en la red dominio, su entorno sirve de predictor sobre las posibles referencias que pueden darse en la descripción de un objeto concreto ligado a ese concepto.

La sintaxis, se utiliza en la medida en la que provee información importante desde el punto de vista semántico, y en particular para agrupar palabras en UNIDADES semánticas, para eliminar posibles ambigüedades, y sobre todo para determinar el ámbito de los cuantificadores permitiendo una imbricación correcta de las subredes generadas durante el análisis.

El analizador es una ATN semántica, ya que lo que se obtiene es la representación (en forma de red) de la consulta, y lleva a cabo únicamente las acciones ligadas a este fin.

Conclusión

Nuestro trabajo (Verdejo 80 a,b,c) se enmarca en un proyecto de construcción de un sistema pregunta-respuesta cuyo lenguaje de consulta es el castellano. El interés del proyecto presenta un doble aspecto:

Teórico - Estudio de la representación del conocimiento y los procesos de comprensión del lenguaje natural.

Práctico - Aplicar el estudio teórico a la problemática de la comunicación con el ordenador en dominios habitualmente tratados en bases de datos.

En esta perspectiva, el analizador cubre una primera etapa, ya que acepta un subconjunto del castellano, que comprende las frases normalmente utilizadas en este tipo de sistemas.

Resaltamos que su comportamiento es independiente del dominio en que se trabaja y está ligado al formalismo de representación (red semántica) con sus mecanismos asociativos y de inferencia.

Actualmente, el sistema recibe la representación del dominio como un dato inicial. Estamos ampliando el subconjunto del lenguaje, para aceptar frases declarativas de forma que la red dominio pueda crearse de forma interactiva en un diálogo usuario-sistema.

Referencias

- G. Camarero & G. Sanz & M. F. Verdejo (79 a)
Representación del conocimiento mediante redes semánticas
CIL 79. Pg 457-466
- G. Camarero & G. Sanz & M. F. Verdejo (79 b)
Resolución de sistemas de ecuaciones con variables en el conjunto de los contornos de un grafo semántico
4 Congreso AEIA, Madrid. Pg. 711-717
- M. F. Verdejo (80 a)
Un analizador semántico para un sistema pregunta respuesta en lenguaje natural. Tesis doctoral. UCM.
- M. F. Verdejo (80 b)
Un sistema de producción para la imbricación de redes semánticas. Revista de Informática y Automática N. 45
- M. F. Verdejo (80 c)
Análisis predictivo y representación del conocimiento
Convención Informática Latina. Barcelona 1980

ACTA

Lugar : Facultad de Filología, Dpto. de Lengua Española, Universidad de Sevilla.

Organizadores : M.Meya y M.A.Pineda

Asistentes : Nicolás Campos (Univ.Ciudad Real), José Cuenca (Univ.Polit.Madrid), J.García Sanz (CC Univ.Compl.Madrid), Montserrat Meya (Siemens, Alemania), M.Angel Pineña (Univ. Sevilla), Josefa Postigo (Univ.Complutense), Joaquin Rangel (Univ.Barcelona), M.S.Ruiz-perez (Univ.Complutense), A.J.Rubio (Univ. Granada), Luis de Sopena (CC Univ.Bilbao), M.Felisa Verdejo (Univ.San Sebastian)

Motivacion : Circular enviada por M.Meya a los asistentes con el fin de coordinar los trabajos existentes en España dentro del área del Procesamiento del Lenguaje Natural.
Meta inicial era el conocimiento mutuo, y la redacción de un proyecto-estudio a presentar a la Fundación Juan March.

Puntos tratados el 15.4.1983:

1. M.Meya presentó sus intenciones/motivaciones al organizar esta reunion, a saber, la necesidad de coordinar las actividades interdisciplinarias dentro del área del Procesamiento del Lenguaje Natural, y conocer los trabajos existentes en España para no repetir investigaciones y esfuerzos.
2. La negativa de la Fundación Juan March, -expresada pocos días antes - de apoyar un proyecto tal de coordinación junto con proyectos concretos a definir por el grupo, fue debatida junto con la posibilidad de cómo plantear un proyecto realista que pudiera ser apoyado desde fuera.
Como única posibilidad viable se comentó la necesidad de tener a mano un estudio en el que se definan funciones y partes concretas dentro de un proyecto, junto con trabajos realizados antes de ir en busca de apoyo.
3. Presentación e historial breve de los asistentes.

4. Discusión acerca de los medios existentes: ordenadores que dispone cada uno, tipo de financiación que se tiene, libertad o no de decisión acerca del tipo de instalación que uno quiere, proyectos a realizar o en marcha, etc...
5. Debate acerca del tratamiento de textos:
 - vaciado de textos (próximo, fidedigno a la fuente)
 - estandarización
 - automatización de diccionarios
 - paquetes y listas estadísticas
 - análisis de correspondencias
 - análisis morfológico
 - almacenamiento (convenciones utilizadas)
6. Especificación de las metas que quieren obtenerse a nivel lexicográfico:
 - crear un grupo de trabajo
 - crear/fijar una norma de estandarización a nivel formal; p.ej.: convenciones para los signos diacríticos y estructuras formales del texto (párrafo, capítulo, diálogo, etc); a nivel físico: soporte en disquette de 5 pulgadas mejor que de 8 pulgadas, qué densidad, cómo etiquetar una cinta, etc...
 - especificar qué léxicos se necesitan para posteriores trabajos, qué categorización tienen los existentes, posibilidad de acumular los léxicos existentes...
 - cómo interpretar el diccionario de la RAE de forma que lo pudieran utilizar todos (cómo generar los asientos prescindiendo de las definiciones)
 - especificación unificada para todos cara a la documentación de los listados y léxicos que se disponen.

Meta final: disponer de una base de datos lingüística bien especificada y documentada, y de un mínimo de herramientas común.

J. García ha traído a Sevilla en cinta su "jeudemo" para poder hacer copias y ponerla a disposición de todos. El paquete de programas "jedemo" es un sistema para el tratamiento de textos que acepta cualquier texto independientemente de las reglas de transcripción usadas. Obtiene índices, concordancias, frecuencias, listados en orden inverso, etc...

Se considera necesario hacer un fichero de los programas existentes para poder intercambiarlos.

Como será la especificación/documentación de los programas queda por fijar.

Puntos tratados el 16.4.1983:

7. Entraron en liza aquellos trabajos que pasan del tratamiento a procesamiento de textos y datos lingüísticos.

L. de Sopena trató el tema interfases en lengua natural a partir de una exposición breve de su trabajo en USL query language en IBM.

Las interfases en lengua natural como medio de comunicarse con el ordenador para solicitar información a un banco de datos- aunque una de las aplicaciones más extendidas/discutidas en el procesamiento del lenguaje - se la intentó encuadrar después con J. Cuenca y F. Verdejo como una aplicación más dentro de los modelos de diálogo, front-end.

Igualmente más que a título de discusión se expuso a título de información aquellas actividades del Procesamiento del LN que caen dentro de las técnicas de la Inteligencia Artificial (IA).

F. Verdejo trató de modo veloz exponer su interés por el lenguaje dentro de sus actividades en la IA.

También J. Cuenca que se dedica fundamentalmente a técnicas de deducción automática y representación del conocimiento, se refirió a la lógica con que se deciden, p.ej. los procesos técnicos (p.ej. para crear un sistema experto para la defensa y control de las inundaciones de los ríos).

8. A. Rubio que se dedica al Procesamiento del Habla explicó el día anterior los procedimientos con que trabajan, síntesis por predicción lineal en un sistema monolocutor.

Tras las exposiciones de tan diversa índole quedó claro la necesidad de conocer las diversas técnicas de análisis del lenguaje para conocer, por un lado, qué implementaciones en sistemas se han realizado en España, y por otro, qué tipo de contribuciones teóricas existen dentro de las áreas: representación del conocimiento, inferencia, técnicas de diálogo, etc..

Se abrió, pues, una perspectiva más amplia de procesar el lenguaje teniendo en cuenta también los mecanismos que subyacen al lenguaje.

9. Se discutió la viabilidad de llevar a cabo en el futuro un proyecto en común en algún grupo.

J. Cuenca insistió en la necesidad de definir el "producto", y discutir propuestas acerca de qué tipo de aplicación.

Se habría de tratar, lógicamente, también : forma de trabajo (en la diáspora), criterios, filosofía, organización técnica, hacia que productos sociales, etc...

Como aplicación concreta se propuso una interfaz en lengua natural para una aplicación bancaria.

Para la tarde se dejaron los temas:

10:utilidad práctica, implicaciones científicas y teóricas, necesidad de crear una revista/boletín, de fundar sociedad, etc...

Por la tarde se acordó:

- M. Ruiperez recogería información para definir los estatutos.
- En Junio se publicaría un boletín/revista con exposiciones de cada uno en una extensión máxima de 6 páginas.
- Como lugar de la próxima reunión se propuso el CC de la Universidad Complutense de Madrid. La fecha quedó abierta.
- Las firmas para formar sociedad se recogerían por correo certificado.
- J. Sanz enviaría la documentación del "Jeudemo" a M.A. Pineda y a L. Sopena. La copia de la cinta se hizo esa misma mañana.

M. Moya

.....

Información recogida

ordenador	Modelo	Sistema operativo	lenguaje	area	Persona	Entidad
BM/270	Optimist 80	varios ?	Basic	lexicografia:indices, concordancias;semantica.	Nicolas Campos	Univ.Ciudad Real
igital	Vax 11-750	V 3.1	Pascal,Interlisp	procesos deductivos en redes inferenciales; sistemas de expertos	J.Cuena	Univ.Polit.Madrid
ata General	Eclipse C/350	AOS	Pascal	fonetica,morfologia,.. procesamiento del habla	F.Casacubieta	Univ.Valencia
-	--	--	--	lexicografia,semantica, banco de datos,indices,..	E.Garcia Camarero	
BM	370/158	OS	Fortran IV	lexicografia:indices,..	G.Garcia Sanz	CC Univ.Complutense
iemens	7760	BS-2000	SPL,Pascal	morfologia,interfases en leng.natural	M.Meya	Siemens AG München
BM	4341	VM/SP-CMS	PL/1	tratamiento de textos, diseno de lenguajes	H.Mielgo	Univ.Barcelona Fac.de Matematicas
igital	PDP-11/60	RSX-11 M	Fortran IV Plus	fonetica acustica: procesamiento del habla	E.Munoz Merino	Esc.Tec.Sup.Ingen. Univ. Madrid
igital	VAX/PC 325	RSX-11M,VSM,P/OS- RSX-11 M	Fortran Plus Basic	fonetica,morfologia, lexicografia:indices,..	M.A.Pineda	Dpto.Lengua Univ.Sevilla
BM/270	Optimist 80	varios	Basic	lexicografia:indices,..	J.Postigo	CC Univ.Complutense
BM	4341	VM/ SP	Cobol,Assembler	lingüística cuantitativa, textual,morfologia,...	J.Rafel	Univ.Barcelona, Dpto.Catalan
ewlett	HP87	-	Extended Basic	tratamiento de textos	M.S.Ruiperez	Univ.Complutense Dpto.Fil.Griega
igital	PDP-11/60	RSX-11 M	Fortran IV Plus	sisntesis del habla	J.Romano	Univ.Tecnica Munich
ata General	Eclipse 250	AOS	Fortran V	procesamiento del habla	A.J.Rubio	Univ.Granada
erkin Elser	3220	OS/32	Fortran	morfologia,lexicografia, interfases en leng.nat.,..	L.de Sopena	Univ.Bilbao
igital	VAX/11/750	VAX/VMS	Basic,Pascal, LISP, Prolog	representacion del cono- cimiento, anal.leng.nat., lenguajes de especificacion	M.F.Verdejo	Univ.San Sebastian Fac. de Informatica
igital	VAX/VMS 11-	V.3.1	Fortran	morfologia,sintaxis, semantica,..:parser,..	J.Vinoly	Intersoftware

CONSTITUCIÓN DE LA SOCIEDAD ESPAÑOLA PARA EL PROCESAMIENTO DEL LENGUAJE NATURAL

ESTATUTOS

TÍTULO I: DE LA CONSTITUCIÓN, DOMICILIO, Y FINES

Artículo 1. Se constituye en Sevilla el 6 de Octubre de 1983 La Sociedad Española para el Procesamiento del Lenguaje Natural, que tendrá ámbito nacional.

Artículo 2. El domicilio social de la sociedad es la Facultad de Informática del País Vasco, San Sebastian. En caso de que la Asamblea General estime conveniente su traslado el nuevo domicilio deberá ser comunicado a todos los socios.

Artículo 3. Son fines de la Sociedad :

- a) Promoción y difusión de la investigación dentro del procesamiento del lenguaje natural.
- b) establecer facilidades para intercambio de información y materiales científicos.
- c) la colaboración e intercambio con instituciones nacionales y extranjeras que cultiven la misma investigación.
- d) la organización y promoción de seminarios, simposios y conferencias sobre el área.
- e) promoción de publicaciones dentro del área.
- f) la creación de juntas técnicas y grupos de trabajo.

TÍTULO II: DE LOS SOCIOS Y DE LA ORGANIZACIÓN DE LA SOCIEDAD ESPAÑOLA PARA EL PROCESAMIENTO DEL LENGUAJE NATURAL.

Artículo 4. Los socios serán numerarios y correspondientes. Sólo los socios numerarios gozarán de los derechos de representación, elegibilidad y elección. Los socios deberán aportar la cuota que la asamblea estipule.

Artículo 5. Puede ser socio numerario todo aquel que desee entrar en la sociedad , y que demuestre interés y dedicación por el área.

Artículo 6. Pueden ser socios correspondientes aquellas personas, entidades o sociedades que ayuden a la Sociedad Española para el Procesamiento del Lenguaje Natural a la realización de sus fines, faciliten su relación con otras entidades semejantes, o cooperen con ella.

Artículo 7. A la condición de socio se puede renunciar por petición propia, o bien ésta será retirada por la Junta Directiva por incumplimiento de las obligaciones inherentes.

Artículo 8. La Junta Directiva estará compuesta por: a) un Presidente, b) un Vicepresidente, c) un Secretario-tesorero, d) tres vocales.

Artículo 9. Será competencia de la Junta Directiva:

- a) representar los intereses de los socios a nivel nacional e internacional
- b) someter la gestión administrativa y organizatoria a la Asamblea General
- c) organizar coloquios, reuniones, seminarios y otras actividades científicas.
- d) publicar y distribuir la revista/boletín órgano de la Sociedad, así como de los estudios y monografías de interés dentro del área del procesamiento del lenguaje natural.

Artículo 10. El Presidente convoca las reuniones de la Junta Directiva, de las sesiones científicas, de la Asamblea General y dirige sus deliberaciones. El Vicepresidente sustituirá al Presidente en caso de ausencia de este. La Asamblea puede igualmente convocar reuniones extraordinarias de la Junta Directiva y organizar actividades científicas.

Artículo 11. El Secretario-tesorero despacha los asuntos generales de la Sociedad, lleva la contabilidad y redacta el presupuesto y la memoria anual.

Artículo 12. La Asamblea General ordinaria de socios se reunirá una vez al año para la aprobación de la gestión de la Junta. La elección de los miembros de la Junta directiva será bianual. La votación en la Asamblea puede tener lugar por carta enviada al Secretario de la Sociedad. La elección de los cargos de la Junta es vigente si se obtiene el 50% de los votos. La votación es solamente válida si recoge como mínimo un tercio de los votos de los miembros.

Artículo 13. La convocatoria de la Asamblea General extraordinaria deberá hacerse por citación individual de los socios con expresión del orden del día.

Artículo 14. Los miembros de la Junta directiva se considerarán cesados en sus cargos, bien por expirar el período de su mandato, bien porque la Asamblea tome un acuerdo que represente desaprobación o censura a su gestión.

Artículo 15. Con una antelación mínima de un mes a la terminación de su mandato la Junta directiva saliente abrirá el plazo de presentación libre de candidaturas a los cargos de Presidente, Vicepresidente, secretario-tesorero y vocales., los cuales serán elegidos por mayoría simple de votos en la Asamblea General.

TÍTULO III: DE LOS MEDIOS Y RECURSOS DE LA SOCIEDAD

Artículo 16. Para alcanzar sus fines la Sociedad Española para el Procesamiento del Lenguaje Natural podrá organizar conferencias, cursillos, publicar revistas o monografías; enviar un representante a organizaciones o actividades paralelas ajenas.

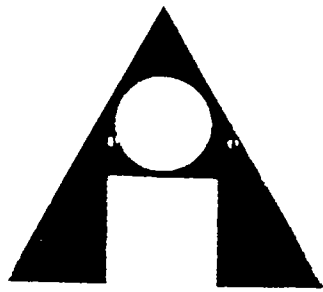
Artículo 17. La Sociedad celebrará reuniones de trabajo a intervalos variables según las necesidades de los socios.

Artículo 18. Los fondos de la Sociedad estarán formados por las cuotas de sus socios y por la venta de sus publicaciones. Los fondos de la Sociedad estarán depositados en un banco a nombre de la misma. Será objeto de reglamentación interna el manejo de estos fondos.

DISPOSICIÓN ADICIONAL

Artículo 19. La Junta Directiva solicitará de la autoridad competente, de acuerdo con la legislación, la declaración de la Sociedad Española para el Procesamiento del Lenguaje Natural como Asociación de Utilidad Pública.

Artículo 20. Las modificaciones del presente Reglamento podrán hacerse: a) por iniciativa de la Junta directiva de la Sociedad; b) por aprobación de la mayoría de los socios.



ijcai-83

En Karlsruhe, Alemania Federal, del 8 al 12 de Agosto de 1983, tuvo lugar el 8º Congreso de la Federación Internacional de Inteligencia Artificial.

Inteligencia Artificial (IA) es como saben muchos una ciencia muy nueva que intenta crear modelos del comportamiento humano de modo que estos puedan ser aplicados a un ordenador, para que éste simule un comportamiento inteligente.

Esta tarea exige esfuerzos enormes y una colaboración interdisciplinaria como no ha existido antes dentro de la ciencia y la técnica. Se trata, pues, de acudir a la psicología, a la lingüística, a las ciencias cognitivas, técnicas de aprendizaje, lógica formal, etc... para crear modelos de cómo actúa un ser inteligente.

En el Congreso se organizaron las sesiones acerca de los siguientes temas:

- programación automática
- modelación cognitiva
- sistemas expertos
- representación del conocimiento
- aprendizaje y adquisición del conocimiento
- programación en lógica
- procesamiento del lenguaje natural
- robótica
- visión

Dentro del área de Procesamiento del lenguaje natural se trataron los siguientes temas:

- gramáticas independientes del contexto (postchomkismo. La conferencia de Gazdar despertó enorme interés por las consecuencias que informáticamente ello acarrea).
- teorías formales de la adquisición del lenguaje
- generación (tema muy concurrido)
- interases en lengua natural y aplicaciones (cabe resaltar las conferencias de P.Martin, D.Appell y F.Pereira con TEAM, y Jaime Carbonell con XCALIBUR)

- parsing, implementación de gramáticas y morfología.
- discurso, diálogo, complejidad narrativa, etc...
- reconocimiento del habla y síntesis (exposiciones muy modestas).

Las sesiones eran paralelas dada la enorme cantidad de conferencias presentadas. Hubo paneles y mesas redondas sobre:

1. Panel on the Fifth Generation Project.
2. Artificial Intelligence: Its impact on human Occupations and Distribution of Income.
3. A panel on AI and Databases

Adjuntamos la referencia de las conferencias presentadas en el area del Procesamiento del Lenguaje Natural.

J. Cuena, M. Meya y F. Verdejo tienen los Proceedings . Si deseáis fotocopia de alguna conferencia podéis dirigiros a cualquiera de nosotros.

M. Meya

The National Conference on Artificial Intelligence

KNOWLEDGE REPRESENTATION AND PROBLEM SOLVING

<i>An Overview of Meta-Level Architecture</i> Michael R. Genesereth, Stanford University	119
<i>Finding All of the Solutions to a Problem</i> David E. Smith, Stanford University	373
<i>Communication and Interaction in Multi-Agent Planning</i> Michael Georgeff, SRI International	125
<i>Data Dependencies on Inequalities</i> Drew McDermott, Yale University	266
<i>KRYPTON: Integrating Terminology and Assertion</i> Ronald J. Brachman and Hector J. Levesque, Fairchild Laboratory for Artificial Intelligence Research; Richard E. Fikes, Xerox Palo Alto Research Center	31
<i>The Denotational Semantics of Horn Clauses as a Production System</i> J.-L. Lassez and M. Maher, University of Melbourne, Australia	229
<i>Theory Resolution: Building in Nonequational Theories</i> Mark E. Stickel, SRI International	391
<i>Improving the Expressiveness of Many Sorted Logic</i> Anthony G. Cohn, University of Warwick, England	84
<i>The Bayesian Basis of Common Sense Medical Diagnosis</i> Eugene Charniak, Brown University	70
<i>Analyzing the Roles of Descriptions and Actions in Open Systems</i> Cari Hewitt and Peter de Jong, Massachusetts Institute of Technology	162
<i>Proving the Correctness of Digital Hardware Designs</i> Harry G. Barrow, Fairchild Laboratory for Artificial Intelligence Research	17
<i>A Chess Program That Chunks</i> Murray Campbell and Hans Berliner, Carnegie-Mellon University	49
<i>The Decomposition of a Large Domain: Reasoning About Machines</i> Craig Stanfill, University of Maryland	387
<i>Reasoning About State From Causation and Time in a Medical Domain</i> William J. Long, Massachusetts Institute of Technology	251
<i>The Use of Qualitative and Quantitative Simulations</i> Reid G. Simmons, Massachusetts Institute of Technology	364

<i>An Automatic Algorithm Designer: An Initial Implementation</i> Elaine Kant and Allen Newell, Carnegie-Mellon University	177
<i>On Inheritance Hierarchies With Exceptions</i> David W. Etherington, University of British Columbia, Canada; Raymond Reiter, University of British Columbia, Canada, and Rutgers University	104
<i>Default Reasoning as Likelihood Reasoning</i> Elaine Rich, University of Texas at Austin	348
<i>Default Reasoning Using Monotonic Logic: A Modest Proposal</i> Jane Terry Nutter, Tulane University	297
<i>A Theorem-Prover for a Decidable Subset of Default Logic</i> Philippe Besnard, Rene Quiniou, and Patrice Quinton, IRISA—INRIA Rennes, Campus de Beaulieu, France	27
<i>Derivational Analogy and Its Role in Problem Solving</i> Jaime G. Carbonell, Carnegie-Mellon University	64

COGNITIVE MODELING

<i>Three Dimensions of Design Development</i> Neil M. Goldman, USC/Information Sciences Institute	130
<i>Six Problems for Story Understanders</i> Peter Norvig, University of California at Berkeley ...	284
<i>Planning and Goal Interaction: The Use of Past Solutions in Present Situations</i> Kristian J. Hammond, Yale University	148
<i>A Model of Learning by Incremental Analogical Reasoning and Debugging</i> Mark H. Burstein, Yale University	45
<i>Modeling Human Knowledge of Routes: Partial Knowledge and Individual Variation</i> Benjamin Kuipers, Tufts University	216
<i>STRATEGIST: A Program That Models Strategy-Driven and Content-Driven Inference Behavior</i> Richard H. Granger, Kurt P. Eiselt, and Jennifer K. Holbrook, University of California at Irvine	139
<i>Learning Operator Semantics by Analogy</i> Sarah A. Douglas, Stanford University and Xerox Palo Alto Research Center; Thomas P. Moran, Xerox Palo Alto Research Center	101
<i>An Analysis of a Welfare Eligibility Determination Interview: A Planning Approach</i> Eswaran Subrahmanian, Carnegie-Mellon University	39

PROCEEDINGS AAAI-83

VISION AND ROBOTICS

<i>A Variational Approach to Edge Detection</i> John Canny, Massachusetts Institute of Technology	54
<i>Surface Constraints From Linear Extents</i> John R. Kender, Columbia University	187
<i>An Iterative Method for Reconstructing Convex Polyhedra From External Gaussian Images</i> James J. Little, University of British Columbia, Canada	247
<i>Two Results Concerning Ambiguity in Shape From Shading</i> Michael J. Brooks, Flinders University of South Australia, Australia	36
<i>Perceptual Organization as a Basis for Visual Recognition</i> David G. Lowe and Thomas O. Binford, Stanford University	255
<i>Model-Based Interpretation of Range Imagery</i> Darwin T. Kuan and Robert J. Drazovich, Advanced Information & Decision Systems	210
<i>A Design Method for Relaxation Labeling Applications</i> Robert A. Hummel, New York University	166
<i>Appropriate Lengths Between Phalanges of Multijointed Fingers for Stable Grasping</i> Tokuji Okada and Takeo Kanade, Carnegie-Mellon University	301
<i>Find-Path for a PUMA-Class Robot</i> Rodney A. Brooks, Massachusetts Institute of Technology	40
<i>Rule Based Strategies for Image Interpretation</i> T. E. Weymouth, J. S. Griffith, A. R. Hanson, and E. M. Riseman, University of Massachusetts	429

NATURAL LANGUAGE

<i>Recursion in TEXT and Its Use in Language Generation</i> Kathleen R. McKeown, Columbia University	270
<i>Repairing Miscommunication: Relaxation in Reference</i> Bradley A. Goodman, Bolt Beranek & Newman, Inc.	134
<i>Tracking User Goals in an Information-Seeking Environment</i> Sandra Carberry, University of Delaware	59
<i>Reasons for Beliefs in Understanding: Applications of Non-Monotonic Dependencies to Story Processing</i> Paul O'Rorke, University of Illinois at Urbana-Champaign	306
<i>RESEARCHER: An Overview</i> Michael Lebowitz, Columbia University	232
<i>Phonotactic and Lexical Constraints in Speech Recognition</i> Daniel P. Huttenlocher and Victor W. Zue, Massachusetts Institute of Technology	172
<i>Deterministic and Bottom-Up Parsing in Prolog</i> Edward P. Stabler, Jr., University of Western Ontario, Canada	383

<i>MCHART: A Flexible, Modular Chart Parsing System</i> Henry Thompson, University of Edinburgh, Scotland	408
<i>Inference-Driven Semantic Analysis</i> Martha Stone Palmer, SDC, A Burroughs Company, and University of Pennsylvania	310
<i>Mapping Between Semantic Representations Using Horn Clauses</i> Ralph M. Weischedel, University of Delaware	424
<i>Constraining a Deterministic Parser</i> J. Bachenko, D. Hindle, and E. Fitzpatrick, Naval Research Laboratory	8
<i>QE-III: A Formal Approach to Natural Language Querying</i> James Clifford, New York University	79
<i>An Overview of the Penman Text Generation System</i> William C. Mann, USC/Information Sciences Institute	261
<i>Interactive Script Instantiation</i> Michael J. Pazzani, The MITRE Corporation	320

LEARNING

<i>Episodic Learning</i> Dennis Kibler and Bruce Porter, University of California at Irvine	191
<i>Human Procedural Skill Acquisition: Theory, Model and Psychological Validation</i> Kurt VanLehn, Xerox Palo Alto Research Center	420
<i>A Production System for Learning Plans From an Expert</i> D. Paul Benjamin and Malcolm C. Harrison, New York University	22
<i>Operator Decomposability: A New Type of Problem Structure</i> Richard E. Korf, Carnegie-Mellon University	206
<i>Schema Selection and Stochastic Inference in Modular Environments</i> Paul Smolensky, University of California at San Diego	378
<i>Why AM and Eurisko Appear to Work</i> Douglas B. Lenat, Stanford University; John Seely Brown, Xerox Palo Alto Research Center	236
<i>Learning Physical Descriptions From Functional Definitions, Examples, and Precedents</i> Patrick H. Winston and Boris Katz, Massachusetts Institute of Technology; Thomas O. Binford and Michael Lowry, Stanford University	433
<i>A Problem-Solver for Making Advice Operational</i> Jack Mostow, USC/Information Sciences Institute	279
<i>Generating Hypotheses to Explain Prediction Failures</i> Steven Salzberg, Yale University	352
<i>Learning by Re-Expressing Concepts for Efficient Recognition</i> Richard M. Keller, Rutgers University	182
<i>Learning: The Construction of A Posteriori Knowledge Structures</i> Paul D. Scott, University of Michigan	359

Sponsored by the American Association for Artificial Intelligence

August 22-26, 1983
Washington, D.C.

RESUMENES DE LOS PROCEEDINGS AAAI.

RESEARCHER: AN OVERVIEW

Michael Lebowitz

Department of Computer Science
Computer Science Building, Columbia University
New York, NY 10027

Abstract

Described in this paper is a computer system, RESEARCHER, being developed at Columbia that reads natural language text in the form of patent abstracts and creates a permanent long-term memory based on concepts generalized from these texts, forming an intelligent information system. This paper is intended to give an overview of RESEARCHER. We will describe briefly the four main areas dealt with in the design of RESEARCHER: 1) knowledge representation, where a canonical scheme for representing physical objects has been developed, 2) memory-based text processing, 3) generalization and generalization-based memory organization that treats concept formation as an integral part of understanding, and 4) generalization-based question answering.

1 Introduction

Natural language processing and memory organization are logical components of intelligent information systems. At Columbia, we are developing a computer system, RESEARCHER, that reads natural language text in the form of patent abstracts (disc drive patents provide the initial domain) and creates a permanent long-term memory based on generalizations that it makes from these texts. In terms of task, RESEARCHER is similar to IPP [Lebowitz 80; Lebowitz 83], a program that read and remembered news stories. The need to deal with complex object representations and descriptions has introduced a whole new range of problems not considered for that program.

PHONOTACTIC AND LEXICAL CONSTRAINTS IN SPEECH RECOGNITION

Daniel P. Huttenlocher and Victor W. Zue
Department of Electrical Engineering and Computer Science
Massachusetts Institute of Technology
Cambridge, Massachusetts 02139

ABSTRACT

We demonstrate a method for partitioning a large lexicon into small equivalence classes, based on sequential phonetic and prosodic constraints. The representation is attractive for speech recognition systems because it allows all but a small number of word candidates to be excluded, using only gross phonetic and prosodic information. The approach is a robust one in that the representation is relatively insensitive to phonetic variability and recognition error.

DETERMINISTIC AND BOTTOM-UP PARSING IN PROLOG

Edward P. Stabler, Jr.

University of Western Ontario
London, Canada

ABSTRACT

It is well known that top-down backtracking context free parsers are easy to write in Prolog, and that these parsers can be extended to give the power of ATN's. This report shows that a number of other familiar parser designs can be very naturally implemented in Prolog. The top-down parsers can easily be constrained to do deterministic parsing of LL(k) languages. Bottom-up backtrack parsers can also be elegantly implemented and similarly constrained to do deterministic LR(k) parsing. Very natural extensions of these LR(k) parser designs suffice for deterministic parsing of natural languages of the sort carried out by the Marcus(1980) parser.

INTERACTIVE SCRIPT INSTANTIATION

Michael J. Pazzani
The MITRE Corporation
Bedford, MA 01730

ABSTRACT

The KNOBS [ENGELMAN 80] planning system is an experimental expert system which assists a user by instantiating a stereotypical solution to his problem. SNUKA, the natural language understanding component of KNOBS, can engage in a dialog with the user to allow him to enter components of a plan or to ask questions about the contents of a database which describes the planning world. User input is processed with respect to several knowledge sources including word definitions, scripts which describe the relationships among the scenes of the problem solution, and four production system rule bases which determine the proper database access for answering questions, infer missing meaning elements, describe how to conduct a conversation, and monitor the topic of the conversation. SNUKA differs from GUS [BOBROW 77], a dialog system similar to SNUKA in its goals, in its use of a script to guide the conversation, interpret indirect answers to questions, determine the referents of nominals, perform inferences to answer the user's questions, and decide upon the order of asking questions of the user to maintain a coherent conversation. SNUKA differs from other script-based language understanders such as SAM [CULLINGFORD 78] and FRUMP [DEJONG 79] in its role as a conversational participant instead of a story understander.

RECURSION IN TEXT AND ITS USE IN LANGUAGE GENERATION

Kathleen R. McKeown
Department of Computer Science
Columbia University
New York, NY 10027

ABSTRACT

In this paper, I show how *textual structure* is recursive in nature; that is, the same rhetorical strategies that are available for constructing the text's macro-structure are available for constructing its sub-sequences as well, resulting in a hierarchically structured text. The recursive formalism presented can be used by a generation system to vary the amount of detail it presents for the same discourse goal in different situations.

1 Introduction

Texts and dialogues often contain embedded units which serve a sub-function of the text or dialogue as a whole. This has been noted both by Grosz [GROSZ 77] in her observations on task dialogues and by Reichman [REICHMAN 81] in analyses of informal conversations. In this paper, I show how *textual structure* is recursive in nature; that is, the same rhetorical strategies that are available for constructing the text's macro-structure are available for constructing its sub-sequences as well, resulting in a hierarchically structured text. This complements Grosz's view of hierarchical text structure as a mirror of hierarchical task structure. A generation system can use recursion to generate a variety of different length texts from a limited number of discourse plans which specify appropriate textual structures. In the following sections, I present a formulation of recursive text structure, an example of its use in the fully implemented TEXT generation system [MCKEOWN 82a], and finally a description of some recent work on the application of this mechanism to automatically generating the appropriate level of detail for a user.

REPAIRING MISCOMMUNICATION: RELAXATION IN REFERENCE

Bradley A. Goodman

Bolt Beranek and Newman Inc.
10 Moulton Street
Cambridge, MA. 02238

ABSTRACT

In natural language interactions, a speaker and listener cannot be assured to have the same beliefs, contexts, backgrounds or goals. This leads to difficulties and mistakes when a listener tries to interpret a speaker's utterance. One principal source of trouble is the description constructed by the speaker to refer to an actual object in the world. The description can be imprecise, confused, ambiguous or overly specific; it might be interpreted under the wrong context. This paper explores the problem of resolving such reference failures in the context of the task of assembling a toy water pump. We are using actual protocols to drive the design of a program that plays the part of an apprentice who must interpret the instructions of an expert and carry them out. A primary means for the apprentice to repair such descriptions is by relaxing parts of the description.

TRACKING USER GOALS IN AN INFORMATION-SEEKING ENVIRONMENT

Sandra Carberry

Department of Computer Science
University of Delaware
Newark, Delaware 19711

ABSTRACT

This paper presents a model for hypothesizing and tracking the changing task-level goals of a speaker during the course of an information-seeking dialogue. It allows a complex set of domain-dependent plans, forming a hierarchical structure of component goals and actions. Our model builds the user's plan as the dialogue progresses, maintains both a local and a global plan context, and differentiates between past goals and goals currently pursued by the user. This research is part of a project to develop a robust natural language interface. If an utterance cannot be interpreted normally or a response cannot be generated due to pragmatic overshoot, the strong expectations about the utterance provided by our context model can be used as an aid in processing the input and producing useful responses.

CONSTRAINING A DETERMINISTIC PARSER

J. Bachenko, D. Hindle, and E. Fitzpatrick
Computer Science and Systems Branch
Information Technology Division
Naval Research Laboratory
Washington, D.C. 20375

ABSTRACT

At the Naval Research Laboratory, we are building a deterministic parser, based on principles proposed by Marcus, that can be used in interpreting military message narrative. A central goal of our project is to make the parser useful for real-time applications by constraining the parser's actions and so enhancing its efficiency. In this paper, we propose that a parser can determine the correct structures for English without looking past the "left corner" of a constituent, i.e. the leftmost element of the constituent along with its lexical category (e.g. N, V, Adj). We show that this Left Corner Constraint, which has been built into our parser, leads quite naturally to a description of verb complements in English that is consistent with the findings of recent linguistic theory, in particular, Chomsky's government and binding (GB) framework.

MCHART: A FLEXIBLE, MODULAR CHART PARSING SYSTEM

HENRY THOMPSON
Department of Artificial Intelligence
University of Edinburgh
Hope Park Square, Meadow Lane
Edinburgh EH8 9NW

ABSTRACT

One of the most attractive properties of the active chart parsing methodology (Kay 1980; Thompson and Ritchie 1983) is the distinction it makes possible between essential bookkeeping mechanisms, scheduling issues, and details of grammatical formalisms. MCHART is a framework within which active chart parsing systems can be constructed. It provides the essential bookkeeping mechanisms, and carefully structured interfaces for the specification of scheduling and grammatical formalism. The resulting flexibility makes it useful both for pedagogical purposes and for quick prototyping. The system is available in UCILISP, FranzLisp, and Interlisp versions, together with a simple lexicon facility, example parsers and detailed documentation.

1. ACTIVE CHART PARSING: A BRIEF OVERVIEW

Constraints on space make it impossible to describe the active chart parsing methodology in any detail. This section merely presents a sketch, to outline the main points and establish salient terminology. For more detail see Kay or Thompson and Ritchie (op. cit.).

Active chart parsing takes the idea of a well-formed substring table, which is a passive lattice consisting of one edge per constituent (either initial or discovered), and makes it active by using edges to record partial or incomplete hypothesized constituents as well. Such active edges contain not only a description of the partial contents of the hypothesized constituents, but also some indication of what more is required to complete them. It follows that the central point in the active chart parsing process occurs when an active and inactive edge meet for the first time. In each such circumstance, one or more new edges may be added if the inactive edge (partially) satisfies the requirements for completion of the active edge. How this is determined will of course depend on the grammatical formalism being used. We call this whole operation the fundamental rule. Note that it speaks of adding new edges, not changing old ones - this is crucial, as will become clear later.

The fundamental rule alone is not of course sufficient to transform an initial chart of lexical edges into a fully analysed one. For this a source of empty active edges, that is,

initial hypotheses about possible constituents, is also required. Such edges typically arise either in response to the addition of inactive edges to the chart (data-driven), or of active edges (hypothesis driven), based on the contents of the edges and the grammar. We call the choice a particular parser makes on this issue the rule invocation strategy, as it determines when and how rules in the grammar enter the chart as hypothesized constituents.

Both the fundamental rule and the rule invocation strategy give rise to non-determinism. The addition of a single edge to the chart, either active or inactive, may provoke a number of applications of the fundamental rule, each in turn possibly provoking new edges. Similarly, a new edge may provoke the rule invocation strategy to hypothesize a number of potential constituents. We call the question of how to order the pursuit of these various possible additions the search strategy.

AN OVERVIEW OF THE PENMAN TEXT GENERATION SYSTEM

William C. Mann
USC Information Sciences Institute
4676 Admiralty Way
Mann del Rey, CA 90291

ABSTRACT

The problem of programming computers to produce natural language explanations and other texts on demand is an active research area in artificial intelligence. In the past, research systems designed for this purpose have been limited by the weakness of their linguistic bases, especially their grammars, and their techniques often cannot be transferred to new knowledge domains.

A new text generation system, Penman, is designed to overcome these problems and produce fluent multiparagraph text in English in response to a goal presented to the system. Penman consists of four major modules: a knowledge acquisition module which can perform domain-specific searches for knowledge relevant to a given communication goal; a text planning module which can organize the relevant information, decide what portion to present, and decide how to lead the reader's attention and knowledge through the content; a sentence generation module based on a large systemic grammar of English; and an evaluation and plan-perturbation module which revises text plans based on evaluation of text produced.

Development of Penman has included implementation of the largest systemic grammar of English in a single notation. A new semantic notation has been added to the systemic framework, and the semantics of nearly the entire grammar has been defined. The semantics is designed to be independent of the system's knowledge notation, so that it is usable with widely differing knowledge representations, including both frame-based and predicate-calculus-based approaches.

INFERENCE-DRIVEN SEMANTIC ANALYSIS

Martha Stone Palmer

SDC - A Burroughs Company
and
University of Pennsylvania
ABSTRACT

A primary problem in the area of natural language processing is the problem of semantic analysis. This involves both formalizing the general and domain-dependent semantic information relevant to the task involved, and developing a uniform method for access to that information. Natural language interfaces are generally also required to have access to the syntactic analysis of a sentence as well as knowledge of the prior discourse to produce a semantic representation adequate for the task. This paper briefly describes previous approaches to semantic analysis, specifically those approaches which can be described as using *templates*, and corresponding multiple levels of representation. It then presents an alternative to the template approach, inference-driven semantic analysis, which can perform the same tasks but without needing as many levels of representation.

MAPPING BETWEEN SEMANTIC REPRESENTATIONS USING HORN CLAUSES¹

Ralph M. Weischedel
Computer & Information Sciences
University of Delaware
Newark, DE 19711

ABSTRACT

Even after an unambiguous semantic interpretation has been computed for a sentence in context, there are at least three reasons that a system may map the semantic representation R into another form S.

1. The terms of R, while reflecting the user view, may require deeper understanding, e.g. may require a version S where metaphors have been analyzed.
2. Transformations of R may be more appropriate for the underlying application system, e.g. S may be a more nearly optimal form. These transformations may not be linguistically motivated.
3. Some transformations may depend on non-structural context.

Design considerations may favor factoring the process into two stages, for reasons of understandability or for easier transportability of the components.

This paper describes the use of Horn clauses for the three classes of transformations listed above. The transformations are part of a system that converts the English description of a software module into a formal specification, i.e. an abstract data type.

REASONS FOR BELIEFS IN UNDERSTANDING: APPLICATIONS OF NON-MONOTONIC DEPENDENCIES TO STORY PROCESSING

Paul O'Rourke

Coordinated Science Laboratory
University of Illinois at Urbana-Champaign
Urbana, IL 61801

ABSTRACT

Many of the inferences and decisions which contribute to understanding involve fallible assumptions. When these assumptions are undermined, computational models of comprehension should respond rationally. This paper crossbreeds AI research on problem solving and understanding to produce a hybrid model ("reasoned understanding"). In particular, the paper shows how non-monotonic dependencies [Doyle79] enable a schema-based story processor to adjust to new information requiring the retraction of assumptions.

QE-III: A FORMAL APPROACH TO NATURAL LANGUAGE QUERYING

James Clifford
Graduate School of Business Administration
New York University

ABSTRACT

In this paper we present an overview of QE-III, a language designed for natural-language querying of historical databases. QE-III is defined formally with a Montague Grammar, extended to provide an interpretation for questions and temporal reference. Moreover, in addition to the traditional syntactic and semantic components, a formal pragmatic interpretation for the sentences of QE-III is also defined.

RECURSION IN TEXT AND ITS USE IN LANGUAGE GENERATION

Kathleen R. McKeown
Department of Computer Science
Columbia University
New York, NY 10027

ABSTRACT

* In this paper, I show how *textual structure* is recursive in nature: that is, the same rhetorical strategies that are available for constructing the text's macro-structure are available for constructing its sub-sequences as well, resulting in a hierarchically structured text. The recursive formalism presented can be used by a generation system to vary the amount of detail it presents for the same discourse goal in different situations.

CALENDARIO

9.-12. Sept. 1983

Int. Conference on Lexicography
LEXeter'83, Exeter (Inglaterra)

19.-23. Sept. 1983

IFIP'83, Paris

28.-30. Sept. 1983

Colloque int. "L'Oeuvre d'Emile
Benveniste: Bilan et perspectives"
Tours

24.-26. Oct. 1983

ACM'83 Conference, New York

19.-21. Oct. 1983

Third ASSP Workshop on Multidimen-
sional Digital Signal Processing,
Stanford Sierra Camp, California

Diciembre- Enero 1983-84

Curso sobre Inteligencia
Artificial, "Artificial Intellig-
ence hands-on introduction,
en BATTELLE, 7, Route de Drize
1227 Carouge (Ginebra)

12.-15. Dic. 1983

CHI'83 Conference Human Factors
in Computing Systems, Boston
Massachusetts.

The theme of the conference is system usability. Individual papers, sessions, and tutorials are invited on all topics relating to the performance and satisfaction of the user of computer systems. In addition empirical studies of the effect of interactive systems on organization behavior are welcomed. Investigations of the impact on the user of various input and output modes are also requested. Reports on human factors studies of intelligent systems and on cognitive models of users are appropriate for presentation. Similarly, work on human factors issues in programming and documentation is requested. Suggestions and recommendations for panels and professional development seminars are urgently solicited (these should be forwarded as soon as possible to the program chair).

Methodology

- ☐ Interface Design Methodology
- ☐ Interface Evaluation Methodology
- ☐ User Cognitive Models
- ☐ Programmer Performance Measurement

Interface and Language Design

- ☐ Human Factors Issues in Object-Based Programming
- ☐ Graphics-Based Interaction
- ☐ Voice Interaction
- ☐ Novel Interface Interaction
- ☐ Novel Language Design
- ☐ Application of Intelligent Systems to Enhance Usability

System Performance Evaluation

- ☐ Usability of Programming Language
- ☐ Critical, Formal Evaluation of Working Systems
- ☐ Organizational Impact of Interactive Systems

30.12-7.1. 1984

1983 IEEE Int. Conference on
Systems, Man and Cybernetics
Bombay and New Delhi

13.-16. Dic. 1983

XIII Simposio de la Soc. Esp.
de Lingüística, Barcelona

2.-6. Abril 1984

ALLC Conference, Computers
in Literary and Linguistic
Research.

Temas de la Conferencia:
Lexicología, Lexicografía,
Mecanización de diccionarios,
Concordancias, Índices, Pro-
cesamiento de textos, Análisis
del contenido, Estilística,
Software and Hardware para
el procesamiento literario y
lingüístico.

1348 Louvain-La-Neuve, Belgica

Contacto: Dr. J. Hamesse

Institut Supérieur de Philosophie,
Chemin d'Aristote, 1

1348 Louvain-la-Neuve

24.-27. Abril 1984

7th European Meeting on Cyber-
netics and Systems Research,
Viena, Austria

March 19-21, 1984

1984 IEEE International Conference on
Acoustics, Speech, and Signal Process-
ing. Location: San Diego, CA (USA). Con-
tact: Dr. Stanley A. White, Rockwell Inter-
national (BB85), P.O. Box 4192, Anaheim,
CA 92803, USA.

CALENDARIO

2.-7. Julio 1984

II Int. Logic Programming Conference
Uppsala, Suecia

Temas:

Applications of Logic Programming
(Nat. Language understanding,
Speech understanding...), etc

5.-10. Agosto 1984

AILA, 7th Int. Congress of Applied
Linguistics, Bruselas

Temas:

language problems in developing
nations, sociolinguistics, psycho-
linguistics, logoco-linguistics...

3.-7. Sept. 1984

ECAI, VI Conferencia Europea de
Inteligencia Artificial

Areas:

Cognitive Modelling, Knowledge
Representation, Natural Language,
Philosophical Issues, etc.

EUSIPCO '83, Second European Signal Processing Conference

Institution: EURASIP

Contact: H.W. Schuessler, Lehrstuhl fuer
Nachrichtentechnik, Universitaet
Erlangen-Nuernberg, Cauerstrasse 7,
D-8520 Erlangen, W. Germany, tel.
(09131) 857100, telex: TFERL 629755

Place: University of Erlangen, Engineering
Department, W. Germany

Date: 12-16 September 1983

Abstract: The aim of the conference is to
cover all aspects of signal processing
theory and practice and to promote the
exchange and crossfertilization of ideas
between individuals working in such a
multi-disciplinary field. Sessions will in-
clude new research results, presentations
of applications, and technological novel-
ties. In addition, tutorial and review
papers will be given. Acknowledged ex-
perts of the field have accepted this task
and will form the main body of the Con-
ference Scientific Committee.

This conference is open to all aspects of
signal processing. The major topics may
be listed in the following form:

- A) Signal processing 1 D
- B) Signal processing 2 D
- C) Speech and sound processing
- D) Detection, estimation, etc.

Advanced Speech Technologies to Revolutionize Interaction of Man and Machine

Talking typewriters, talking computers,
even talking automobiles are moving from
the realm of science fiction to science fact,
according to a report issued today by
Datapro Research Corporation.

The report, 'All About Speech
Technology', assesses the current state-of-
the-art in systems that 'talk' or respond to
voice commands. The new report is design-
ed for businesses that want to explore the
use of voice input and output machines to
streamline communication and improve effi-
ciency in a wide range of operations, from
inventory control to banking and manufac-
turing. Complete technical and applications
data is included on 71 products that utilize
speech technology. These products are sup-
plied by 40 vendors.

The Datapro report examines four dif-
ferent categories of speech technology
equipment: voice response, voice recogni-
tion, voice synthesis and voice store-and-
forward systems, also called voice mail
systems. In addition to information on
specific products, the report provides detail-
ed guidelines to aid companies in evaluating
the capabilities, limitations and costs of cur-
rent systems and their applications in
specific business environments.

Among the many systems that incor-
porate speech response or recognition
capability, the report cites order entry
systems that process oral purchase orders
and devices that allow engineers to verbally
program the operation of machine tools. At
least one Detroit automaker, the report in-
dicates, is expected to introduce a talking
automobile in the 1983 model year. The car
will tell passengers to fasten unbuckled seat
belts, instead of just flashing a light or soun-
ding a warning buzzer.

'All About Speech Technology' is avail-
able for \$19 per copy. To order, contact
Datapro Research Corporation, 1805 Under-
wood Blvd., Delran, NJ 08075 (Datapro
Research Corporation Press Release).

NOTAS

CALL FOR PROPOSALS FOR COLLABORATIVE RESEARCH WORK
IN THE FIELD OF ARTIFICIAL INTELLIGENCE AND PATTERN RECOGNITION

Objective

The Commission of the European Communities, acting in the framework of the pluriannual programme of the Community in data processing has decided to grant financial support to cooperative research work in the field of Artificial Intelligence and Pattern Recognition.

The assistance is not intended to support the cost of the basic research activity carried out by a research team, but rather to support the additional costs of collaboration between teams in different Community countries: travel, subsistence when working abroad, joint workshops, usage of network services and so on. In exceptional cases a more substantial contribution could also be given to projects of particular interest.

Topics

Among the various fields of prospective interest preference will be given initially to those proposals which cover one of the following:

Domains:

- Image processing and understanding
- Speech processing and understanding.

Techniques:

- New algorithms
- Parallel processing
- Expert systems

Applications:

- Robotics
- Office automation
- Medicine
- Remote sensing.

Submission of proposals

Research teams or individuals working in these areas are invited to present proposals for cooperative international research.

Proposals for financial assistance should state

- the objectives and the programme of work for which support is sought and its relationship to the state of the art in the area.

.../...

- the resources (financial and staff) which will be made available for the work with some indication on the way they are to be organized.
- the ways and means of the planned cooperation and the expected benefits to be gained through international cooperation.
- the possible links with other work.
- the aid requested and its justification.
- the timetable for the work and for the payment of aid.
- any other information considered relevant for the appraisal of the proposal.

Should the detailed definition of the proposal require expenditure on travel that cannot be met by the proposer from national resources, a preliminary proposal may be made to cover a contribution to these costs.

The proposals should be addressed before 1 January 1983 to:

The Commission of the European Communities
 Directorate-General III
 Attention Mr J. Desfosses
 200, rue de la loi
 B - 1049 Brussels, Belgium

In return for its aid, the Commission may require progress reports and a commitment to publish the results.

The selection of the proposals will be made with the assistance of the CREST Subcommittee on "Informatics and Information Technologies".

The Commission reserves its right to refuse any proposal without any justification.

BOOKREVIEWS

COMPUTATIONAL GRAMMAR

Graeme D. Ritchie

Sussex: The Harvester Press, 1980.

This book was developed from the author's doctoral dissertation, in which he wrote a program that conducted natural language conversations in order to study the impact that computational linguistics has on developing descriptive grammar. Ritchie does much to show how difficult it is to develop a proper grammar, and thus gives an accurate reflection of the problems of computational linguistics by describing the state of descriptive linguistics and his application of it.

Most of the first quarter of the book is a review of what Ritchie refers to as "frameworks" of linguistic description. He correctly points out that there is, at present, no true theory of either grammar or computational grammar. His survey treats the major schools of current grammatical thought, including Chomsky and Montague. The next segment of the text describes Structural Combining Rules, which are the heart of Ritchie's system. The last long segment is a review of implementation problems. This latter section is unclear both because of its episodic nature and because the reader is not given a feel for why the problems arose. It is in this section that the gap between reader and implementor is not bridged.

One thing that is clear in reading the last section is that even though linguistics and computational linguists have been professing considerable interest and confidence in transformational grammar,

there is still no true implementation of it. The phrase "generative grammar" implies that ability to use a grammar to create (generate) new sentences that would be grammatically correct. This goal has yet to be accomplished, and Ritchie's book gives a hint why: transformational grammar fails to provide a simple, clear solution, the tendency is to quickly create an ad hoc solution based on phrase structure grammars and give up the development of purely transformational or generative programs, preferring to have ones that work (mostly) now. The eventual implementation of the transformational grammars is supposed to be based on what is learned from these less-than-perfect prototypes.

The strengths of the book are the development of the Structural Combining Rules and the fact that the author is aware of the book's limitations. The SCR's are intended as an interface between syntax and semantics, and as such as basically successful.

The weaknesses of the book are the weaknesses of linguistics and computational linguistics in general: it is just one more imperfect approach in a universe of imperfect approaches. The author, in the process of explaining some of his grammar's deficiencies, makes statements about the goals of artificial intelligence research that are incorrect. He says that attempts to use English as the equivalent of a high-level programming language are "misguided." If this is meant to mean that English will never be a medium for programming computers, then he is denigrating the long-range potential of his own research. In discussing the delay time involved in "conversing" with the computer programmed with his natural language interpreter, Ritchie says that processing time should not be a serious consideration. For "natural language" to be truly "natural" excessive delays between conversation exchanges must be avoided.

Tioga Publishing Company announces the publication of a new book entitled *Machine Learning: An Artificial Intelligence Approach*, edited by Ryszard S. Michalski, Jaime G. Carbonell, and Tom M. Mitchell. The book is a collection of chapters by several leading scientists on several aspects of learning by machines, and is intended as a supplementary textbook for general courses in artificial intelligence and also as a primary text for more specialized courses. Specific topics include knowledge acquisition for expert systems, procedure and strategy learning, learning by analogy, learning from advice, learning from examples, learning from observation, modeling human learning, discovery systems, and conceptual data analysis. Introductory and overview chapters and a glossary of basic terms are included for nonspecialists. An extensive bibliography of the field is also included. [Approximately 600 pages, including illustrations, preface, index, glossary, and bibliography. Hardcover, 6 x 9 format. Price \$39.50. ISBN 0-935382-05-4. It is available through William Kaufmann, Inc. 95 First Street, Los Altos, CA 94022.]

READINGS IN ARTIFICIAL INTELLIGENCE

Edited by Bonnie Lynn Webber and Nils J. Nilsson

Thirty-one classic papers on search, deduction, planning, expert systems, and other topics are gathered in this convenient volume for class and reference use. Hard-to-find papers by AI experts are thus made readily accessible for present and future specialists. Representing, as they do, some of the best thinking in the field, they will retain their fascination and interest.

Published by Tioga Publishing Company.
557 pages. Illustrated. Appendix. Notes.
Bibliographies. Index. Paperbound.
\$25.00. ISBN 0-935382-03-8

G

LIBROS

CONCEPTUAL STRUCTURES: INFORMATION PROCESSING IN MIND AND MACHINE

John F. Sowa,
IBM Systems Research Institute

Conceptual Structures, primarily a textbook, looks at knowledge based systems that think for themselves. Combining material from cognitive psychology, linguistics and artificial intelligence, this book is suitable for use in Artificial Intelligence and Advanced Cognitive Psychology courses.

The unifying theme running through the book is the theory of conceptual graphs. These graphs are networks of concepts and conceptual relations that have been developed over the past twenty years as semantic representation in artificial intelligence, linguistics, and cognitive psychology. Under various names, such as semantic networks, correlational nets, conceptual dependency graphs, and cognitive nets, they have been implemented in dozens of computer systems. Although essentially every feature described in the theory has been implemented in some form in some system, no system has it all. This book is the first full-length treatment that synthesizes all of the approaches with a common notation. By presenting all of the features in this single notation, the book helps reduce the bewildering confusion, makes the techniques easier to learn, and leads to more complete, systematic implementations.

CONTENTS

Philosophical Basis • Psychological Evidence • Conceptual Graphs • Reasoning and Computation • Language • Knowledge Engineering • Limits of Conceptualization • Appendix A: Mathematical Background • Appendix B: Conceptual Catalog • Bibliography

Publication Date: March 1983, 300 pp.

Hardbound ISBN 0-201-14472-7

En : American Journal of Computational Linguistics,
Volume 8, Number 2, April-June 1982

ACL Book Series

A new book series has been announced by the Cambridge University Press. It is entitled *Studies in Natural Language Processing*, and has Aravind K. Joshi as the General Editor.

The Association for Computational Linguistics is sponsoring this series of monographs, texts, and collections of papers on scientific and scholarly topics which represent the range of professional activity of concern to its membership. Methods of investigation include the algorithmic, mathematical, and heuristic methods of formal linguistics and artificial intelligence, and the experimental methods of cognitive psychology and psycholinguistics.

It is anticipated that books to appear in the series will be of three general types: research monographs, textbooks, and collections of papers. Research monographs, aimed at the professional reader, will be detailed reports of significant research projects or literature reviews of a size beyond the compass of a journal

article. Appropriately revised doctoral dissertations of particularly high quality may be considered for inclusion in the series in this category.

Courses in computational linguistics are becoming increasingly common in departments of linguistics and computer science, and there are so far few, if any, suitable textbooks. Students must rely largely on scattered journal articles and instructors' lectures. Proposals for texts by authors of established reputations in the field will be entertained. A text might be in the nature of a general survey, or might concentrate on one of the major subareas of research.

The third category of books will be collections of papers. These will in no case be proceedings of general meetings or otherwise heterogeneous collections. Rather, each will be a thematically coherent set of research reports bearing on a single, well-defined topic of current significance, assembled under strong direction and control by an editor or editors who will be responsible for the quality of all the contents, and who will provide a unifying introduction. Such a collection might result from a symposium on a particular topic, to which the participants submit written papers. Or chapters on various aspects of important problems in a particular area may be especially commissioned by a volume editor. Collections of already-published papers, in cases where the subject is a major and significant one, the papers represent definitive treatments of particular approaches to it, and the papers are in scattered and perhaps not readily available sources, may occasionally be considered.

The Editorial Board for the series consists of: Jonathan Allen, *M.I.T.*, Jaime Carbonell, *Carnegie-Mellon University*, Eugene Charniak, *Brown University*, Paul Chapin, *N.S.F.*, Phil Cohen, *Fairchild, Inc.*, Martha Evens, *Illinois Institute of Technology*, Joyce Friedman, *University of Michigan*, Maurice Gross, *University of Paris*, Barbara Grosz, *SRI International*, George Heidorn, *IBM Research*, Ronald Kaplan, *Xerox Palo Alto Research Center*, Hans Karlgran, *KVAL, Sweden*, Lauri Karttunen, *University of Texas, Austin*, Martin Kay, *Xerox Palo Alto Research Center*, Jan Lansbergen, *Philips Research Laboratories, Eindhoven*, Wendy Lehnert, *University of Massachusetts, Amherst*, William Mann, *USC/Information Sciences Institute*, Mitch Marcus, *Bell Laboratories*, Makoto Nagao, *Kyoto University*, Stanley Petrick, *IBM Research*, Ray Perrault, *University of Toronto*, Ellen Prince, *University of Pennsylvania*, Chris Riesbeck, *Yale University*, Jane Robinson, *SRI International*, Naomi Sager, *New York University*, Peter Sgall, *Charles University, Prague*, Candy Sidner, *Bolt Beranek and Newman*, Norman Sondheim, *USC/Information Sciences Institute*, Fred Thompson, *California Institute of Technology*, Walter von Hahn, *University of Hamburg*, Donald Walker, *SRI International*, Bonnie Webber, *University of Pennsylvania*,

LIBROS

LANGUAGE AS A COGNITIVE PROCESS

Volume 1: Syntax

Terry Winograd
Stanford University

This textbook introduces a computational approach to the structure of human language. It presents a number of major linguistic theories within a cognitive computational framework. The fundamental ideas of computation are presented as they apply to language processing, and all of the prominent current approaches to syntax are explained within the framework (including transformational grammar, augmented transition networks, systemic grammar, case grammar, functional grammars and generalized phrase structure grammars). The ideas are developed step-by-step, so students with no previous linguistic background can learn all of the relevant ideas through the text and exercises. There is also a survey of the existing computer systems for parsing natural language and an outline of English syntax.

CONTENTS

Viewing Language as a Knowledge-based Process • Word Patterns and Word Classes • Context-free Grammars and Parsing • Transformational Grammar • Augmented Transition Network Grammars • Feature and Function Grammars • Computer Systems for Natural Language Parsing • Appendices • Bibliography • Indices
1982, 654 pp., illus.
Hardbound ISBN 0-201-08571-2

SUBJECTIVE UNDERSTANDING, COMPUTER MODELS OF BELIEF SYSTEMS

Jaime G. Carbonell
UMI Research Press
Ann Arbor, Michigan
Copyright 1981

This is a very interesting and very important book. It is based on Carbonell's 1979 Yale Ph.D. thesis working under Schank and Abelson and others. The book is only 285 pages long but it has a lot of outstanding material in it. It has a very useful discussion of the whole question of subjective understanding. It appears that there are operating programs running on the DEC 10 or DEC 20 though no program listing are given. It would have been helpful if there had been a little more technical information on the program languages used and the amount of memory and machine capacity used.

In the book there are three main computer models as follows:

1. POLITICS - A model of ideological reasoning.
2. TRIAD - Techniques for analyzing social and political conflicts.
3. MICS - Mixed-initiative conversational system or in other words, computer programs making it possible to converse with humans in fairly normal speech.

Each of these models are explained in good detail with good diagrams and illustrations. Some of the important concepts emerging from these models include, goals trees, goal-based counter-planning, structures of plan conflicts and goal trees representing personality traits.

PLANNING AND UNDERSTANDING

A Computational Approach to Human Reasoning
Robert Wilensky
University of California, Berkeley

This book, representing original research in the areas of Artificial Intelligence and Natural Language Processing, brings together two areas of cognition: common sense problem solving, and natural language understanding. To do so, three theories are presented: a theory of plan generation, a theory of plan-based understanding, and a theory of the structure of plans that underlies both theories of processing. This research has led to the development of several computer programs, including PAM (Plan Applier Mechanism), a story understanding system, and PANDORA (Plan Analysis with Dynamic Organization, Revision, and Application).

CONTENTS

Introduction • Tenets of a Theory of Plans • Planning in Everyday Situations • Meta-Planning • Explanation-Driven Understanding • Goal Relationships • Negative Goal Relationships • Reasoning about Goal Conflict • Reasoning about Goal Competition • Positive Goal Relationships • Computer Implementation - Representation of Task Networks • Computer Implementation - Programs • Bibliography • Indices

Publication Date: November 1982
1983, 167 pp., illus.
Hardbound ISBN 0901-09590-4

The book is clearly written with good tables and specific illustrations of the point being made. It has a number of clear and well-written flow charts describing the processes.

The author is aware of some of the "problems" in the programs though it would appear that these are not really problems as much as indications of future work that needs to be done. Some of the "problems" lead to interesting and amusing examples, one is given on page 183 on the Triad System where a discussion occurs as follows:

Interpretation of U.S. position on politics with a liberal viewpoint. The input is "Russia massed troops on the Czech border." The question is "What should the U.S. do about it?" The computer responded with the "The United States should congratulate Brezhnev."

At first blush this is really quite humorous but the author goes on to explain how this could easily come from the model. In this case, the goal conflict model suggests that we might want to improve diplomatic relations and one way to improve diplomatic relations is to support the actions of the other country. Examples like this show keen insight by the author into some significant implications of his system.

There are a couple of unsettling things as one studies the book, one problem is that it is not at all clear how the author develops the goal trees. It would be nice if he could give the reader some help on how to go about developing new goal trees in a particular situation. It would be nice if these programs could be tied to the "BELIEVER" of Schmidt and Sridharan.

Dr. Gary Carlson
Computer Translation, Inc.
1455 South State Street
Orem, Utah 84057

BOOKREVIEWS