

Análisis morfológico automático del español a través de generación

Alexander Gelbukh^{1,2}, Grigori Sidorov¹, Francisco Velásquez¹

¹ Centro de Investigación en Computación (CIC),
Instituto Politécnico Nacional (IPN),
Av. Juan de Dios Bátiz, esq. con Miguel Othón de Mendizábal,
México D. F., Zacatenco, CP 07738, México
{gelbukh, sidorov}@cic.ipn.mx, fcastillov@ipn.mx

² Universidad Chung-Ang, Seúl, Corea

Resumen

La mayoría de los sistemas de análisis morfológico están basados en el modelo conocido como la morfología de dos niveles. Sin embargo, este modelo no es muy adecuado para lenguajes con alternaciones irregulares de raíz (por ejemplo, el español o el ruso). En este trabajo describimos un sistema computacional de análisis morfológico para el lenguaje español basado en otro modelo, cuya idea principal es el análisis a través de generación. El modelo consiste en un conjunto de reglas para obtener todas las raíces de una forma de palabra para cada lexema, su almacenamiento en el diccionario, la producción de todas las hipótesis posibles durante el análisis y su comprobación a través de la generación morfológica. Se usó un diccionario de 40,000 lemas, a través del cual se pueden analizar más de 2,500,000 formas gramaticales posibles. Para el tratamiento de palabras desconocidas se está desarrollando un algoritmo basado en heurísticas. El sistema desarrollado está disponible sin costo para el uso académico.

Palabras clave: análisis morfológico automático, generación morfológica automática, análisis a través de generación, español.

Abstract

Most widely spread method used in automatic morphological analysis systems is based on the well-known two-level morphology model. However, this model is not well suit for languages with irregular stem alternations (e.g., Spanish or Russian). In this work we describe a system with automatic morphological analysis for Spanish based on another model, based on analysis through generation. This model consists of a set of rules that allow for obtaining of all possible stems for each lexeme, storing them in the dictionary, producing all possible hypotheses during analysis, and verification of the hypotheses through morphological generation. We used a dictionary containing 40,000 lemmas, with which more than 2.500.000 possible grammatical forms can be recognized. For treatment of unknown words, a heuristic-based algorithm is being developed. The system is freely available for academic use.

Key words: automatic morphological analysis, automatic morphological generation, analysis through generation, Spanish.

1. Introducción

La morfología estudia la estructura de las palabras y su relación con las categorías gramáticas del lenguaje. El objetivo del análisis morfológico automático es llevar a cabo una clasificación morfológica de una forma de palabra. Por ejemplo, el análisis de la forma *gatos* resulta en

gato+Noun+Masc+Pl,

que nos indica que se trata de un sustantivo plural con género masculino y que su forma normalizada (lema) es *gato*.

Los lenguajes según sus características morfológicas –básicamente, según la tendencia en la manera de la combinación de morfemas– se clasifican en *aglutinativos* y *flexivos*.

Se dice que un lenguaje es *aglutinativo* si:

- Cada morfema expresa un sólo valor¹ de una categoría gramatical,
- No existen alternaciones de raíces o las alternaciones cumplen con las reglas morfonológicas que no dependen de la raíz específica, como, por ejemplo, armonía de vocales, etc.,
- Los morfemas se concatenan sin alteraciones,
- La raíz existe como palabra sin concatenarse con morfemas adicionales algunos.

Unos ejemplos de los lenguajes aglutinativos son las lenguas turcas (el turco, kazakh, kirguiz, etc.) o el húngaro.

Por otro lado, un lenguaje es *flexivo* si:

- Cada morfema puede expresar varios valores de las categorías gramaticales. Por ejemplo, el morfema *-mos* en español expresa cumulativamente los valores de las categorías *persona* (tercera) y *número* (plural),
- Alternaciones de raíces no son previsibles –sin saber las propiedades de la raíz específica no se puede decir qué tipo de alternación se presentará,
- Los morfemas pueden concatenarse con ciertos procesos morfonológicos no estándares en la juntura de morfemas,
- La raíz no existe como palabra sin morfemas adicionales (por ejemplo, *escrib-* no existe como palabra sin *-ir*, *-iste*, *-ía*, etc.).

¹ Por ejemplo, nominativo es un valor de la categoría gramatical caso.

Unos ejemplos de los lenguajes flexivos son lenguas eslavas (ruso, checo, ucraniano, etc.) o románicas (latín, portugués, español, etc.).

Normalmente, esta clasificación de lenguajes refleja las tendencias; es decir, muy raras veces un lenguaje es absolutamente aglutinativo o flexivo. Por ejemplo, el finlandés es un lenguaje básicamente aglutinativo, aunque con algunos rasgos de un lenguaje flexivo –por ejemplo, varios valores de las categorías gramaticales pueden unirse en el mismo morfema.

En este artículo, sólo consideraremos los lenguajes flexivos, específicamente el español.

En teoría, ya que la morfología de cualquier lenguaje flexivo es finita, cualquier método de análisis basado en diccionario da resultados igualmente correctos. Sin embargo, no todos los métodos de análisis automático son igualmente convenientes en el uso y fáciles de implementar.

En este trabajo presentamos una implementación para el español de un modelo de análisis morfológico automático basado en la metodología de análisis a través de generación. Esta metodología permitió realizar el desarrollo del sistema con un esfuerzo mínimo y aplicar el modelo gramatical del español que se presenta en las gramáticas tradicionales –el más simple y intuitivo.

En el resto del artículo se describen los modelos existentes del análisis morfológico automático y especialmente el modelo de análisis a través de generación, después se presenta el proceso de generación y análisis que se usó en el sistema desarrollado para el español, se explica el procedimiento de preparación de los datos, se presenta brevemente la implementación del sistema y finalmente se dan las conclusiones.

2. Modelos de análisis morfológico automático

En la implementación de los analizadores automáticos morfológicos es importante distinguir:

- Modelo de análisis (el procedimiento de análisis),
- Modelo de gramática que se usa en el analizador (las clases gramaticales de palabras),
- Implementación computacional (el formalismo usado).

La razón de la diversidad de los modelos de análisis es que los lenguajes diferentes tienen la estructura

morfológica diferente, entonces, los métodos apropiados para lenguajes, digamos, con morfología pobre (como el inglés) o para los lenguajes aglutinativos no son los mejores para los lenguajes flexivos como el español o, digamos, el ruso.

La complejidad del sistema morfológico de un lenguaje para la tarea de análisis automático no depende tanto del número de las clases gramaticales ni de la homonimia de las flexiones, sino del número y tipo de las alternaciones en raíces, las cuales no se puede saber sin consultar el diccionario, por ejemplo: *mover*–*nuevo* vs. *dormir*–*durmió* vs. *correr*–*corro*. En este caso, el tipo de alternación es una característica de la raíz, y el único modo de obtener esta información es consultar el diccionario el cual contiene estos datos de antemano. Al contrario, para el caso de algún idioma aglutinativo como el finlandés, existen alternaciones de raíces, pero normalmente son predecibles sin el uso del diccionario.

Un ejemplo de clasificación de los métodos de análisis morfológico es la clasificación propuesta por R. Hausser [1999a, 1999b], donde los métodos se clasifican en los basados en formas, en morfemas y en alomorfos. Para distinguir entre los sistemas basados en alomorfos y los sistemas basados en morfemas también se usa el concepto de «procesamiento de raíces estático vs. dinámico». De hecho, en el método de alomorfos se guardan todos los alomorfos de cada raíz en el diccionario, lo que es el procesamiento estático, mientras que en el método basado en morfemas es necesario generar los alomorfos de un morfema dinámicamente durante el procesamiento.

Consideremos esos métodos con más detalle. Como un extremo, se puede almacenar todas las formas gramaticales en un diccionario, junto con su lema y toda la información gramatical asociada con la forma. Éste es método basado en formas de palabras. Con esta aproximación, un sistema morfológico es sólo una gran base de datos con una estructura simple. Este método se puede aplicar para los lenguajes flexivos aunque no para los aglutinativos donde se puede concatenar los morfemas casi infinitamente. Los computadores modernos tienen la posibilidad de almacenar bases de datos con toda la información gramatical para grandes diccionarios de lenguajes flexivos –un aproximado de 20 a 50 megabytes para el español o el ruso. Para el español un ejemplo de implementación del dicho método es el proyecto MACO+ [Atseria *et al.*, 1998].

Sin embargo, tales modelos tienen sus desventajas, por ejemplo, no permiten el procesamiento de palabras desconocidas. Otra desventaja es la dificultad de agregar las palabras nuevas al diccionario –hay

que agregar cada forma manualmente. Sin embargo, hacer este tipo de trabajo manualmente es muy costoso. Por ejemplo, los verbos españoles tienen por lo menos 60 formas diferentes (sin contar formas con enclíticos). Para evitar este tipo de trabajo manual se tiene que desarrollar los algoritmos de generación, lo que de hecho puede ser una parte significativa de desarrollo de los algoritmos de análisis, como se presenta en este artículo.

Otra consideración a favor del desarrollo de los algoritmos en lugar de usar una base de datos de las formas gramaticales es el punto de vista según el cual los algoritmos de análisis son un método de compresión del diccionario. Esos métodos permiten la compresión en por lo menos más de 10 veces. En nuestros experimentos hicimos la compresión de los diccionarios del ruso y del español en forma de una base de datos con una utilita de compresión estándar (zip). El fichero del resultado para el ruso fue como 30 veces más grande que el diccionario de los sistemas de análisis; en el caso del español, la diferencia entre el tamaño de los ficheros fue como 10 veces a favor del diccionario para el algoritmo.

Una razón más para el uso de los algoritmos de análisis es que tal análisis es necesario para cualquier tarea que involucra el conocimiento morfológico, por ejemplo, para la tarea de dividir palabras con guiones. También, digamos, para la traducción automática o recuperación de información a veces es mejor tener la información de morfemas que forman la palabra y no únicamente la información a nivel de la palabra completa.

Otro tipo de sistemas se basan en almacenamiento de morfemas (de algún alomorfo que se considera el básico) que representan las raíces en el diccionario. Es decir, el diccionario tiene un sólo alomorfo que representa cada morfema. Los demás alomorfos se construyen en el proceso de análisis. El modelo más conocido de este tipo es PC-KIMMO.

Muchos procesadores morfológicos están basados en el modelo de dos niveles de Kimmo Koskenniemi [1983]. Originalmente, el modelo fue desarrollado para el lenguaje finlandés, después se le hicieron algunas modificaciones para diferentes lenguajes (inglés, árabe, etc.). Poco después de la publicación de la disertación de Koskenniemi donde se propuso el modelo, L. Karttunen y otras personas desarrollaron una implementación en LISP del modelo de dos niveles y lo llamaron PC-KIMMO.

La idea básica del modelo KIMMO es establecer la correspondencia entre el nivel profundo, donde se encuentran solamente los morfemas, y el nivel superficial, donde hay alomorfos. Eso explica porque

en este modelo es indispensable construir los alomorfos dinámicamente.

Además, otra idea detrás del modelo PC-KIMMO es enfocarse en el formalismo –los autómatas finitos (transductores), y de tal modo no pensar en implementación de los algoritmos [Beesley and Karttunen, 2003]. Es decir, los transductores por sí es una implementación del algoritmo. La complejidad del modelo gramatical no se toma en cuenta. No obstante, es necesario desarrollar las reglas para poder construir una raíz básica –el alomorfo que se encuentra en el diccionario– de cualquier otro alomorfo. Si en el idioma no existe muy compleja estructura de alternaciones de raíces, entonces sí se puede usar el modelo KIMMO, como se muestra, por ejemplo, en [Karttunen, 2003]. Para los idiomas donde hay muchas alternaciones de raíces no predecibles –como, por ejemplo, el ruso [Bider and Bolehakov, 1976]– es posible aplicar este modelo, pero el desarrollo de los algoritmos resulta mucho más difícil, aunque no imposible porque todo el sistema es finito. Cabe mencionar que la complejidad de los algoritmos de este tipo según algunas estimaciones es NP-completa [Hausser, 1999b: 255].

Sin embargo, las ideas relacionadas con la implementación no deben considerarse como predominantes. Claro, si las demás condiciones son iguales, es preferible tener una implementación ya hecha (como, por ejemplo, un transductor), pero consideramos que la complejidad de los algoritmos es el costo demasiado elevado. Nuestra experiencia muestra que cualquier otro modo de implementación –interpretador de las tablas gramaticales o programación directa de las reglas– es igualmente efectivo tanto en el desarrollo como en el funcionamiento.

Hay otros modelos de análisis morfológico basados en morfemas. Para el español, Moreno y Goñi [1995] proponen un modelo para el tratamiento completo de la flexión de verbos, sustantivos y adjetivos. Este modelo –GRAMPAL– está basado en la unificación de características y depende de un léxico de alomorfos tanto para las raíces como para las flexiones. Las formas de palabras son construidas por la concatenación de alomorfos, por medio de características contextuales especiales. Se hace uso de las gramáticas de cláusulas definidas (DCG) modeladas en la mayoría de las implementaciones en Prolog. Sin embargo, según los autores, el modelo no es computacionalmente eficiente, es decir, el análisis es lento.

La desventaja común de los modelos basados en el procesamiento dinámico de las raíces es la necesidad de desarrollar los algoritmos de construcción de la raíz básica (la que se encuentra en el diccionario) y

aplicarles muchas veces durante el procesamiento, lo que afecta la velocidad del sistema. Por ejemplo, en el español para cualquier *i* en la raíz de la palabra se tiene que probar cambiarla por *e*, por la posible alternación de raíz tipo *pedir* vs. *pido*. Esa regla puede aplicarse muchas veces durante el análisis para posibles variantes de la raíz (como, por ejemplo, *pido*, *pid-*, *pi-*, etc.).

Los algoritmos de construcción de la raíz básica a partir de las otras raíces no son muy intuitivos, por ejemplo, normalmente, no se encuentran en las gramáticas tradicionales de los idiomas correspondientes. Al contrario, las clasificaciones de alternaciones que están en las gramáticas se usan la raíz básica y de esta raíz se construyen las demás raíces. Además, los algoritmos de construcción de la raíz básica son bastante complejos: por ejemplo, para el ruso el número de las reglas se estima alrededor de 1000 [Malkovsky, 1985].

Para evitar la construcción y aplicación de este tipo de algoritmos, se usa el enfoque estático (basado en alomorfos) del desarrollo de sistemas, cuando todos los alomorfos de raíz están en el diccionario con la información del tipo de raíz. Este tipo de sistemas son más simples para el desarrollo (véase la sección 3).

Para construir el diccionario de tal sistema hay que aplicar el algoritmo de la generación de las raíces a partir de la raíz básica. Este procedimiento se hace una sola vez. Los algoritmos de la generación de las raíces a partir de la primera raíz son más claros intuitivamente, se encuentran normalmente en las gramáticas de idiomas y se basan en el conocimiento exacto del tipo de raíz.

El único coste del almacenamiento de todos los alomorfos con la información correspondiente en el diccionario es el aumento del tamaño del mismo, el cual no es muy significativo. En el peor caso –si todas las raíces tuvieran alternaciones– el aumento correspondiente sería al doble (o triple si todas las raíces tienen 3 alomorfos, etc.). Sin embargo, las raíces con alternaciones usualmente son como máximo 20 por ciento de todas las palabras y además la mayoría de las raíces sólo tiene 2 alomorfos, entonces el aumento del tamaño de diccionario es relativamente pequeño. Además, se puede aplicar algún método de compresión del diccionario reduciendo así los datos repetidos [Gel'bukh, 1992].

Otra consideración importante para los modelos de análisis basados en diccionarios de morfemas o alomorfos, es el tipo de modelos gramaticales que se usan.

La solución directa es crear una clase gramatical para cada tipo de alternación de raíz posible, junto con el conjunto de flexiones que caracterizan esta raíz. De tal modo, cada tipo de palabras tiene su propia clase gramatical. Sin embargo, el problema es que tales clases orientadas al análisis no tienen correspondencia alguna en la intuición de los hablantes, y además su número es muy elevado. Por ejemplo, para el ruso son alrededor de 1000 clases [Gel'bukh, 1992], para el checo son alrededor de 1500 clases [Sedlacek and Smrz, 2001].

Otra posible solución es el uso de los modelos de las gramáticas tradicionales. Las gramáticas tradicionales son orientadas a la generación y clasifican las palabras según sus posibilidades de aceptar un conjunto determinado de flexiones (su paradigma). La clasificación según las posibles alternaciones de las raíces se da aparte porque estas clasificaciones son independientes. De tal modo, las clases corresponden muy bien a la intuición de los hablantes y su número es el mínimo posible. Por ejemplo, en la clasificación según los paradigmas para el ruso con su morfología bastante compleja, hay alrededor de 40 clases, para el español se aplica el modelo estándar de las tres clases para los verbos (con las finales *-ar*, *-er*, *-ir*) y una para sustantivos y adjetivos; las diferencias en las flexiones dependen completamente de la forma fonética de la raíz, las peculiaridades adicionales como, por ejemplo, *pluralia tantum*, se dan aparte. Esos modelos son intuitivamente claros y normalmente ya son disponibles –se encuentran en las gramáticas y diccionarios existentes.

Sin embargo, no es muy cómodo usar la clasificación orientada a generación, para el análisis directo. Para eso proponemos usar una metodología conocida como «análisis a través de generación». Nuestra idea es tratar de sustituir el procedimiento de análisis con el procedimiento de generación. Es bien conocido que análisis es mucho más complejo que generación; por ejemplo, podemos comparar los logros de generación de voz con los de reconocimiento de voz.

3. Modelo de análisis a través de generación

Como hemos mencionado, un aspecto crucial en el desarrollo de un sistema de análisis morfológico automático es el tratamiento de las raíces alternas regulares (*deduc-ir* – *deduzc-o*). El procesamiento explícito de tales variantes en el algoritmo es posible pero requiere del desarrollo de muchos modelos y algoritmos adicionales, que no son intuitivamente claros ni fáciles de desarrollar. Para la solución de esta problemática, el sistema que se describe a continuación implementa el modelo desarrollado en

[Gel'bukh, 1992, Sidorov, 1996] y generalizado en [Gelbukh y Sidorov, 2002]. Este modelo consiste en la preparación de las hipótesis durante el análisis y su verificación usando un conjunto de reglas de generación. Las ventajas del modelo de análisis a través de la generación son la simplicidad –no hay que desarrollar algoritmos de construcción de raíces, ni los modelos gramaticales especiales– y la facilidad de implementación. El sistema desarrollado para el español se llama AGME (Analizador y Generador de la Morfología del Español).

Como hemos mencionado, hay dos asuntos principales en el desarrollo del modelo para análisis morfológico automático:

- 1) cómo tratar los alomorfos –estática o dinámicamente, y
- 2) qué tipo de modelos gramaticales se usa.

La idea básica del modelo propuesto de análisis a través de generación es guardar los alomorfos de cada morfema en el diccionario –procesamiento estático, que permite evitar el desarrollo de complejos algoritmos de transformaciones de raíces inevitable en el procesamiento dinámico– y usar los modelos de las gramáticas tradicionales intuitivamente claros.

3.1 Proceso de Generación

El proceso de la generación se desarrolla de la siguiente manera. Tiene como entrada los valores gramaticales de la forma deseada y la cadena que identifique la palabra (cualquiera de las posibles raíces o el lema).

- Se extrae la información necesaria del diccionario;
- Se escoge el número de la raíz necesaria según las plantillas (véase la Sección 4.2);
- Se busca la raíz necesaria en el diccionario;
- Se elige la flexión correcta según el algoritmo desarrollado. El algoritmo es bastante simple y obvio: por ejemplo, para el verbo de la clase 1 en primera persona, plural, indicativo, presente, la flexión es *-amos*, etc.,
- La flexión se concatena con la raíz.

3.2 Proceso de Análisis

El proceso general de análisis morfológico usado en nuestra aplicación es bastante simple: dependiendo de la forma de palabra de entrada, se formula algu-

na(s) hipótesis de acuerdo con la información del diccionario y la flexión posible, después se generan las formas correspondientes para tal(es) hipótesis, si el resultado de generación es igual a la forma de entrada, entonces la hipótesis es correcta. Por ejemplo, para la flexión *-amos* y la información del diccionario para la raíz que corresponde al verbo de la clase 1, se genera la hipótesis de primera persona, plural, indicativo presente (entre otras), etc.

Las formas generadas según las hipótesis se comparan con la original, en caso de coincidencia las hipótesis son correctas.

Más detalladamente, dada una cadena de letras (forma de palabra), para su análisis se ejecutan los siguientes pasos:

1. Quitar letra por letra sus últimas letras, así formulando la hipótesis sobre el posible pinto de división entre la raíz y la flexión (también siempre se verifica la hipótesis de la flexión vacía «» = $\emptyset$).
2. Verificar si existe la flexión elegida. Si no existe la flexión, regresar al paso 1.
3. Si existe la flexión, entonces hallar del diccionario la información sobre la raíz y llenar la estructura de datos correspondiente; si no existe la raíz, regresar al paso 1. En este momento no verificamos la compatibilidad de la raíz y la flexión –esto se hace en generación.
4. Formular la hipótesis.
5. Generar la forma gramatical correspondiente de acuerdo a la hipótesis y la información del diccionario.
6. Si el resultado obtenido coincide con la forma de entrada, entonces la hipótesis se acepta. Si no, el proceso se repite desde el paso 3 con otra raíz homónima (si la hay) o desde el paso 1 con otra hipótesis sobre la flexión.

Nótese que es importante la generación porque de otro modo algunas formas incorrectas serían aceptadas por el sistema, por ejemplo, **acuerdamos* (en lugar de *acordamos*). En este caso existe la flexión *-amos* y la raíz *acuerd-*, pero son incompatibles, lo que se verifica a través de la generación.

En caso del español es necesario procesar los enclíticos. Se ejecuta un paso adicional antes de empezar el proceso de análisis –los enclíticos se especifican en el programa como una lista (*-me*, *-se*, *-selo*, *-melo*, etc.). Siempre se verifica la hipótesis de haber un enclítico al final de la cadena.

4. Modelos usados

En el español, los procesos flexivos ocurren principalmente en los nombres (sustantivos y adjetivos) y verbos. Las demás categorías gramaticales (adverbios, conjunciones, preposiciones, etc.), presentan poca o nula alteración flexiva. El tratamiento de estas últimas se realiza mediante la consulta directa al diccionario.

3.1. Morfología nominal

La variedad de designaciones a que aluden los dos géneros y la arbitrariedad en muchos casos de la asignación del género (masculino o femenino) a los sustantivos impiden determinar con exactitud lo que significa realmente el género. Es preferible considerarlo como un rasgo que clasifica los sustantivos en dos categorías diferentes, sin que los términos *masculino* o *femenino* prejuzguen ningún tipo de sentido concreto [Llorac, 2000].

No existen reglas estándar para la flexión de género en sustantivos. Por lo tanto en nuestro programa se almacenan todas las formas de sustantivos singulares en el diccionario (por ejemplo, *gato* y *gata*). Los adjetivos siempre tienen ambos géneros, entonces sólo una raíz se almacena en el diccionario: por ejemplo, *bonit-* para tanto *bonito* como *bonita*. Ahora bien, el tratamiento de la flexión del número puede ser modelado mediante un conjunto de reglas, así que es suficiente tener la única clase gramatical, porque las reglas dependen de la forma fonética de la raíz.

Por ejemplo, las formas nominales que se terminan en una consonante que no sea /s/, agregan *-es* en su pluralización (por ejemplo, *árbol* – *árboles*). Por otra parte, los nombres que se terminan en vocal *-á*, *-í*, *-ó*, *-ú* tienden a presentar un doble plural en *-s* y *-es* (*esquí*, *esquíes*; *tabú*, *tabúes*; la información de doble plural se da con una marca en el diccionario), aunque algunos de ellos sólo admiten *-s* (*mamás*, *papás*, *dominós*, etc.). La información sobre el plural no estándar se representa a través de las marcas en el diccionario para las raíces correspondientes.

3.2. Morfología verbal

Clasificamos a los verbos en regulares (no presentan variación de raíz, como *cantar*), semiirregulares (no más de cuatro alomorfos de raíces, como *buscar*) e irregulares (más de cuatro variantes de raíz, como *ser*, *estar*).

Afortunadamente, la mayoría de los verbos en español (85%) son regulares. Para éstos, usamos los tres modelos de conjugación tradicionales (representa-

dos, por ejemplo, por los verbos *cantar*, *correr* y *partir*).

Se usan doce modelos de conjugación verbal diferentes para los verbos semiirregulares según sus alternaciones de la raíces. Nótese que esta clasificación es independiente de los modelos de paradigmas. Cada modelo tiene su tipo de alternación y su plantilla de raíces. Por ejemplo, en el modelo A1 se encuentra el verbo *buscar* (entre otros). Tiene dos raíces posibles, en este caso *busc-*, *busqu-*; la segunda raíz se usa para el presente de subjuntivo, primera persona del singular del pretérito indefinido de indicativo y en algunos casos del imperativo; la primera raíz se usa en todos los demás modos y personas.

Se usa una plantilla (cadena de números) para cada modelo de conjugación semiirregular. Cada posición representa una conjugación posible; por ejemplo, la primera posición representa la primera persona del singular del presente de indicativo, las últimas posiciones hacen referencia a las formas no personales. Los números usados en la plantilla son de 0 a 4, donde 0 indica que no existe la forma correspondiente; 1 indica el uso de la raíz original; 2, 3 y 4 son las otras raíces posibles. Por ejemplo, para el modelo A1 se usa la siguiente plantilla:

```
111111111111211111111111111111222222111111
11111111111111221
```

Esta estructura nos facilita el proceso de generación de las formas verbales. Nótese que son 61 posibles formas, ya que no tomamos en cuenta las formas verbales compuestas (como, por ejemplo, *haber buscado*) porque cada su parte se procesa por separado.

Al ser mínimo el número de los verbos completamente irregulares (como *ser*, *estar*, *haber*), su tratamiento consistió en almacenar todas sus formas posibles en el diccionario. El proceso de análisis para estas palabras consiste en generar la hipótesis de un verbo irregular con la flexión vacía; esta hipótesis se verifica a través de la generación, la cual en este caso consiste en buscar la palabra en el diccionario, obtener todas sus variantes y desplegar el campo de la información.

4. Preparación de los Datos

La preparación preliminar de datos consiste de los siguientes pasos principales:

- Describir y clasificar todas las palabras del lenguaje (español) en las clases gramaticales y las marcas adicionales, como, por ejemplo, plu-

ralia tantum. Esta información se toma completamente de los diccionarios existentes;

- Convertir la información léxica disponible, en un diccionario de las raíces. Sólo la primera raíz tiene que ser generada en este paso;
- Aplicar los algoritmos de generación de raíces para generar todas las raíces, copiándoles la información de la primera raíz y asignando el número de la raíz generada.

Se diseñó una estructura de almacenamiento de datos como se muestra en la tabla 1. Para los datos mostrados, el campo *Palabra* contiene el lema, el campo *Base* contiene la raíz, el campo *Info* contiene la clase gramatical, los campos *Marca1*, *Marca2* contienen las marcas gramaticales adicionales. Por ejemplo, el campo *Marca1* del registro 2 (*P*) indica que se trata de una *pluralia tantum* (es decir, al tratar de generar su singular, obtendríamos un error) y para los últimos dos registros indica el modelo de conjugación semiirregular al que pertenece el verbo. El campo *Marca2* para los últimos dos registros indica la raíz original (1) y la segunda raíz posible (2).

Tabla 1. Estructura del diccionario de raíces.

Base	Palabra	Info	Marca1	Marca2
gato	gato	N		
gafa	gafas	N	P	
acert-	acertar	VI	M1	1
aciert-	acertar	VI	M1	2

La información para la preparación de este diccionario se tomó de los diccionarios existentes explicativos y bilingües.

5. Implementación

La base de datos (el diccionario) es una tabla de Paradox donde se almacenan las raíces y otra información de las mismas, como se mostró en la sección 4.

El sistema se desarrolló en C++. Cuenta con una interfaz que permite escoger la forma gramatical para generar o para introducir la palabra a analizar. También existe una versión del sistema que lee y procesa un fichero grande de texto.

6. Conclusiones

Se presentó un sistema para el análisis morfológico, que implementa el modelo de comprobación de hipótesis a través de generación. Las ventajas de este modelo de análisis reflejadas en su implementación son su simplicidad y claridad, lo que resultó en muy reducido tiempo de implementación: el desarrollo de los algoritmos principales sólo tomó varios días.

El diccionario actual tiene un tamaño considerable: 40,000 lemas, incluyendo 23,400 sustantivos, 7,600 verbos y 9,000 adjetivos.

Es importante mencionar que el sistema AGME no sobregenera ni sobreanaliza, es decir, sólo se procesan las formas correctas.

Se está trabajando sobre la forma de tratar las palabras desconocidas (una aproximación inicial es la de formular una heurística del más parecido). Como trabajo futuro se sugiere considerar los procesos de derivación (*bella* $\Rightarrow$ *belleza*) y composición (*agua* + *fiesta* $\Rightarrow$ *aguafiestas*).

El sistema está disponible como un fichero EXE o DLL de Windows, sin costo alguno para el uso académico.

Agradecimientos

Este trabajo fue realizado con el apoyo parcial del gobierno de México (CONACyT, SNI), IPN (CGEPI-IPN, COFAA, PIFI) y RITOS-2 del Subprograma VII de CYTED. El primer autor actualmente se encuentra realizando una estancia sabática en la Universidad Chung-Ang.

Referencias

- [Atseria *et. al.*, 1998] Atserias, J., J. Carmona, I. Castellón, S. Cervell, M. Civit, L. Márquez, M.A. Martí, L. Padró, R. Placer, H. Rodriguez, M. Taulé, J. Turmó Morphosyntactic analysis and parcing of unrestricted Spanish texts. Proc of LREC-98, 1998.
- [Beesley and Karttunen, 2003] Beesley, K. B. and L. Karttunen. Finite state morphology. CSLI publications, Palo Alto, CA, 2003.
- [Bider and. Bolshakov, 1976] Bider, I. G. and I. A. Bolshakov. Formalization of the morphologic component of the Meaning - Text Model. 1. Basic concepts (in Russian with a separate translation to English). ENG. CYBER. R., No. 6, 1976, p. 42-57.
- [Gel'bukh, 1992] Gel'bukh, A.F. Effective implementation of morphology model for an inflectional natural language. *J. Automatic Documentation and Mathematical Linguistics*, Allerton Press, vol. 26, N 1, 1992, pp. 22-31.
- [Gelbukh and Sidorov, 2002] Gelbukh, A. and G. Sidorov. Morphological Analysis of Inflective Languages through Generation. *Revista Procesamiento de lenguaje natural*, España, vol. 29, 2002, pp 105-112. (2002)
- [González and Vigil, 1999] González, B. M. y C. Ll. Vigil. *Los Verbos Españoles*. 3ª Edición, España, Ediciones Colegio de España. 1999. 258 p. (1999)
- [Hausser, 1999a] Hausser, R. *Three Principled Methods of Automatic Word Form Recognition*. Proc. of VEXTAL: Venecia per il Tratamento Automatico delle Lingue. Venice, Italy, 1999. pp. 91-100. (1999).
- [Hausser, 1999b] Hausser, Roland. *Foundations of Computational linguistics*. Springer, 1999, 534 p.
- [Karttunen, 2003] Karttunen, L., Computing with realizational morphology. In: A.Gelbukh (ed.) Proc. of CICLing-2003 (4th International Conference on Intelligent Text Processing and Computational Linguistics), February 15-22, 2003, Mexico City. Lecture Notes in Computer Science N 2588, Springer-Verlag, pp. 203-214.
- [Koskenniemi, 1983] Koskenniemi, K. *Two-Level Morphology: A General Computational Model for Word-Form Recognition and Production*. Thesis Doctoral. Universidad de Helsinki. 1983. 160 p. (1983).
- [Llorac, 2000] Llorac, E. *Gramática de la Lengua Española*. España, Ed. Espasa, 2000. 406 p. (2000).
- [Malkovsky, 1985] Malkovsky, M. G. *Dialogue with an artificial intelligence system* (in Russian). Moscow State University, Moscow, Russia, 1985, 213 pp.
- [Moreno and Goñi, 1995] Moreno, A. and J. Goñi. *GRAMPAL: A Morphological Processor for Spanish Implemented in PROLOG*. En: Mar Sessa y María Alpuente, editores, Proceedings of the Joint Conference on Declarative Programming (GULP-PRODE'95), pp. 321-331, Marina di Vietri (Italia), 1995. (1995)
- [Santana *et al.*, 1999] Santana, O., J. Pérez, *et al.* *FLANOM: Flexionador y Lematizador Automático de Formas Nominales*. Universidad de las Palmas de Gran Canaria. Lingüística Española

Actual XXI, 2. Ed. Arco/Libros, S.L. España, 1999. (1999)

[Sedlacek Smrz, 2001] Sedlacek R. and P. Smrz, A new Czech morphological analyzer AJKA. Proc. of *TSD-2001*. LNCS 2166, Springer, 2001, pp. 100–107.

[Sidorov, 1996] Sidorov, G. O. Lemmatization in automatized system for compilation of personal style dictionaries of literature writers (in Russian). In: *Word by Dostoyevsky* (in Russian),

Moscow, Russia, Russian Academy of Sciences, 1996, pp. 266–300.

[Sproat, 1992] Sproat, R. *Morphology and computation*. Cambridge, MA, MIT Press, 1992, 313 p.