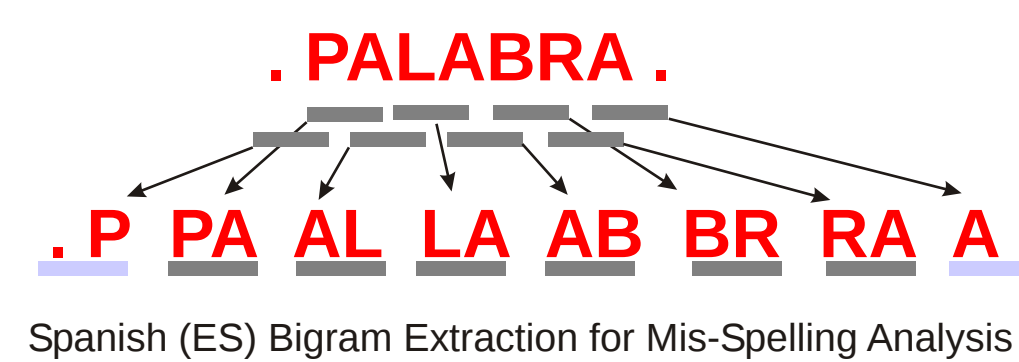


Procesamiento de Lenguaje Natural Robusto



Distribution of Bigrams
(spanish, 48k root-words)

La distribución de bigramas y trigramas es una característica que se repite de forma similar en casi todos los idiomas. Por la forma en que se presenta, permite ser usada para determinar no solo el idioma sino para restaurar errores.

Segmentación y Lematización

La separación en sílabas, palabras y símbolos de un texto no es una tarea trivial, más aún si se desean averiguar los lemas a la vez de tolerar errores, y dar alternativas ortográficas. Normalmente esto es realizado mediante autómatas finitos (*llamados tokenizadores o levers*) basados en lenguajes del tipo *expresión regular*, para reconocer y aislar las secuencias de caracteres que forman las palabras, síglas y números.

Para lograr esto, se ha portado *Jlex* de Java a **C#** bajo **.NET 2.0**, creando una herramienta para Visual Studio.

La lematización con corrección ortográfica, en idiomas flexivos como el español es un problema **NP-duro**, debido a la combinación posible de aijos, amenada por los errores. La contribución es lograr esto en forma efectiva yeficiente.

El análisis gramatical y semántico suele ser un problema complejo cuando el número de terminales (*tokens*) es alto. Nuestro modelo para Español posee más de 700, dada la gran cantidad de conjugaciones verbales y flexiones posibles. Hemos modelizado la gramática española con un parser de tipo GLR el cual hace un buen análisis sin perder generalidad ni velocidad, pese a ser NP-duro.

Se han creado 'utilitarios' académicos para poder trabajar con los diferentes algoritmos de similitud existentes, compararlos entre sí, modelar sus múltiples factores y generar luego las planillas de excel para realizar el trabajo estadístico y gráficas.

Compilador para motor de Diálogo Natural (DDL & C# > MSIL)
 Modelización de procesos cognitivos para diálogo (Dialog Objects)
 Parsing Condicional (Enriquecido con ontologías y contexto)
 Compilador/Parser/Chunker NLP/GLR+ (Con Schrödinger Tokens)
 Modelos Cognitivos (Matemáticas, Conjuntos, Lógica & Unidades)
 Ambiente de Ejecución Cognitivo (DDL Dialog Objects Runtime)
 Comprensión Artificial Superficial por Inferencia Morfológica.
 Resolver Deixis, Co-Referencia y Anáforas (yo/mi/mío, tú/su/suyo..)

andres.hohendahl@fi.uba.ar jfzelasco@fi.uba.ar

Desafíos

En computación, el ingreso de datos y los errores son un problema grave, para lidiar con esto hay que vencer numerosos escollos. Se debe de reconocer y etiquetar ese texto ‘*sucio*’, estimando el idioma y sabiendo si ‘eso’ es pronunciable o si es solamente ruido del tipo ‘*un gato caminando en el teclado*’, o si son siglas, etc. Deben identificarse las palabras, locuciones y términos, números, horas y fechas, siglas, acrónimos y abreviaturas, unidades, monedas, fórmulas matemáticas y otros fenómenos escritos.

Sistemas Usados

Procesamiento de afijos: Algoritmo de análisis de Prefijos, Sufijos e Infijos combinados basados en estándares y diccionarios de libre disponibilidad (*Open Office/ASpell/IsPELL*)

Identificación Estadística de Idioma ⁽¹⁾: Provee medidas para poder reconocer el idioma y si las palabras son pronunciables.

Corrección de Ortografía y Restaurador Léxico: Usamos algoritmos estadísticos con heurística para identificar palabras erradas y brindar un ‘paquete’ de palabras posibles, incluyendo una

Características Obtenidas

- ✓ Módulo analizador de multi-etiquetado robusto con corrector ortográfico, aceptando múltiples idiomas en una misma instancia.

✓ Etiqueta palabras, símbolos y locuciones mediante un estándar internacional, expandido semanticamente basado en *EAGLES 2.0*⁽³⁾ incluyendo una salida adicional en norma *Penn TreeBank* (**en**).

- ✓ Reconoce números en formatos enteros, científicos y hasta hexadecimales, incluyendo Números Romanos como: *MCMXIV*

✓ Identifica Números dichos con palabras como ‘dos mil ciento treinta con 30 céntimos’ tanto en inglés como en español. Adicionalmente cuando etiqueta cardinales, numerales y ordinales, provee el cardinal del cual provienen, muy útil para comprensión artificial basado en semántica, luego del procesamiento gramatical.

✓ Acepta en español e inglés, 925 unidades (SI, CGS, MKS, etc.)
~260 países ~1100 monedas (ISO 4127), ~6000 locuciones (es),
fecha/hora bajo múltiples formatos y normas (ISO 8601, etc.)

- ✓ Controlado por aprox. 65 parámetros especializados.

✓ Analiza la expansión de ‘jerga’ donde hay palabras mezcladas con números, en mensajes cortos y chat; por ej: **salu2 = saludos**

App: Enseñar a Escribir

VAHIEMA → BALLENA (d~0.69 dT: 1 ms)
VAYEMA → BALLENA (d~0.29 dT: 1 ms)

La Similitud Fonética tiene un costado inesperado: permite diseñar algoritmos eficientes para corregir palabras mal escritas por personas que recién están aprendiendo la relación entre los sonidos y las letras. Esto se basa en la "proximidad" fonética con palabras reales del diccionario. Esto combinado con sistemas pedagógicos, constituye un poderoso sistema de Enseñanza y Asistencia Pedagógica en donde las computadoras ya tendrán razgos mas humanos y afables, reduciendo la brecha cultural y digital

La reparación de formas mal escritas o corrección de ortografía, normalmente se basa en permutación y sustitución de letras, pero es difícil saber cual de todas esas permutaciones conviene hacer antes, pues el número de permutaciones y sustituciones es un problema NP-duro. Hemos avanzado creando algoritmos de predicción '*blandos*' que ponen el ojo en donde es más probable que haya un error, similar a lo que un '*humano*' siente intuitivamente.

Referencias

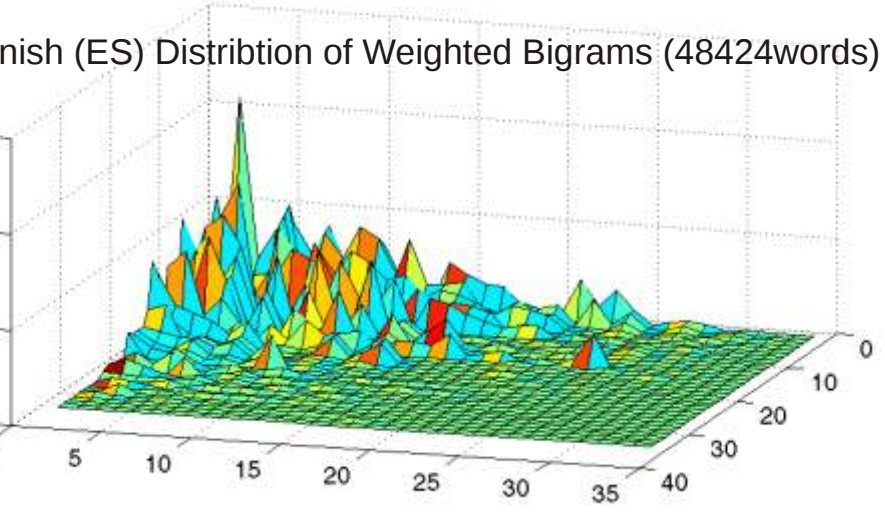
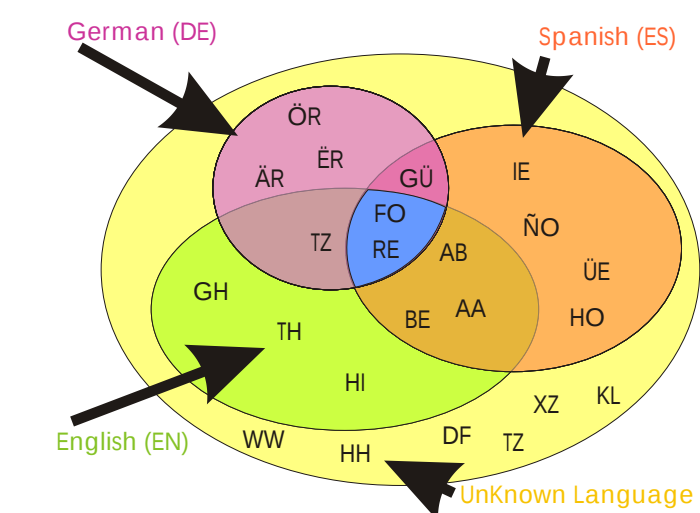
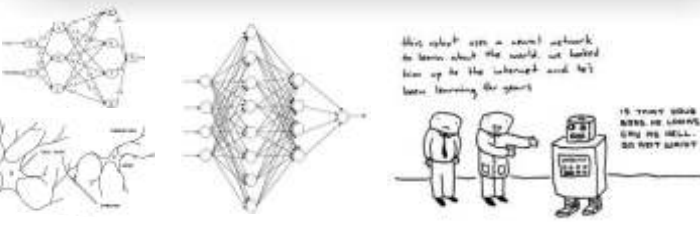
(1) A. Hohendahl, J. Zelasco "Efficient algorithm for fast language detection applied to Natural Language Processing" WICC2006 - ISBN 950-9474-35-5 Art 694. Univ. Morón. Bs.As. Arg.

(2) Hohendahl, A.T.; Zanutto, B. S.; Wainelboim, A. J.; "Development of a Phonetic Similarity Algorithm for written text" SLAN2007. X Latinoamerican Congress of Neuro-Psychology 2007

(3) Our Language detection (1) provides useful information to determine if it is worth trying to correct misspelling or if the input is in some foreign language or is simply garbage.

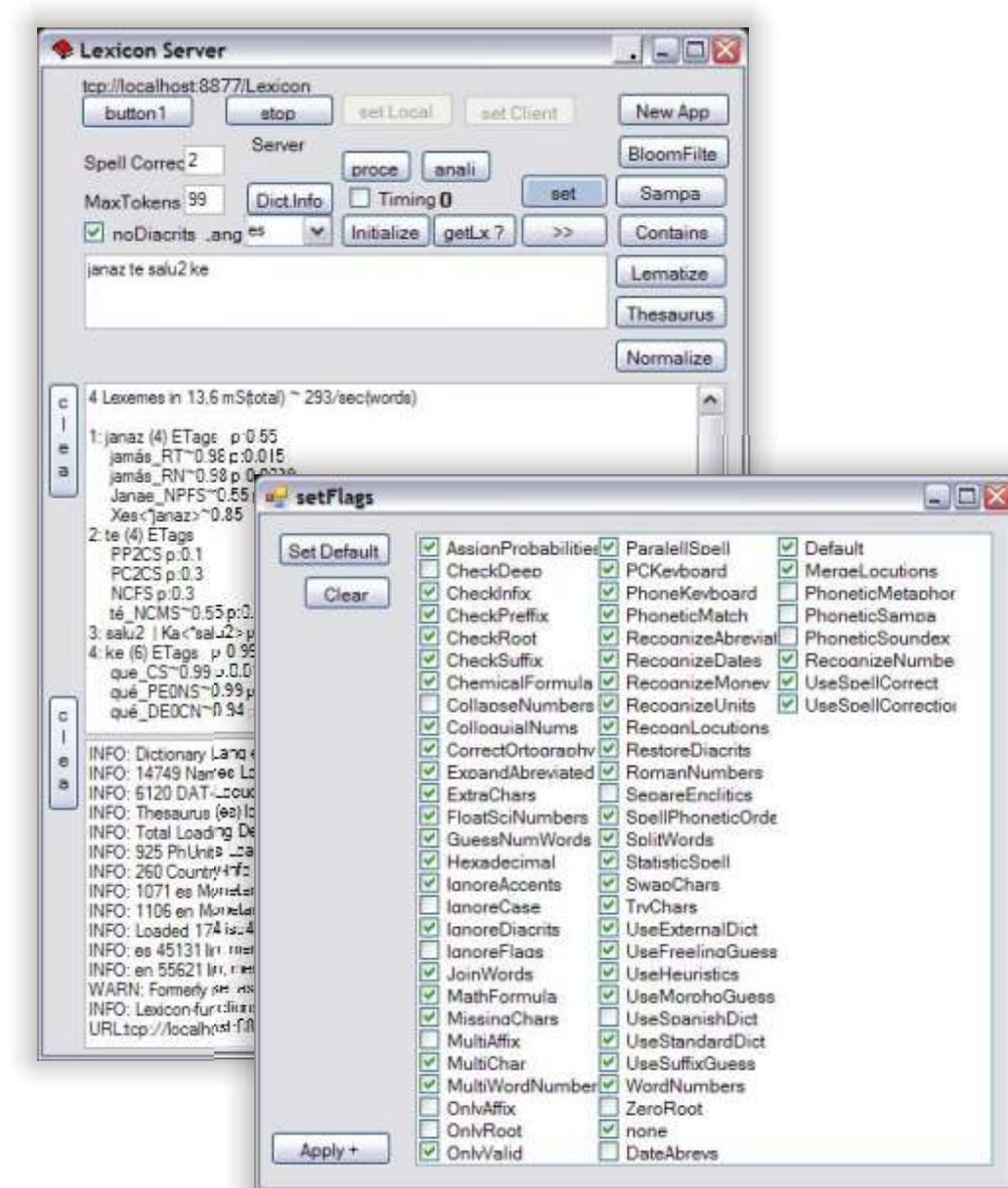
(4) **EAGLES 2.0** / Parole enhanced by means of adding infix semantic classification and with combination of more than 200 specialized tags, useful for text understanding.

(5) Desktop System: XP Sp3 .NET 2.0 2Gb RAM 2.6Ghz i386 E5200 (using single processor core)



La distribución estadística de bigramas entre varios idiomas, provee valiosa información para estimar el idioma de una palabra o frase desconocida.

App.: Servidor Léxico



El servicio creado permite distribuir los recursos de un sistema de procesamiento de lengua, sistema mediante un servidor remoto quien mantiene actualizadas sus bases de datos automáticamente en un solo lugar, proveyendo de funciones lógicas a todo módulo que lo solicite, mediante protocolos de alta performance, evitando la duplicación de recursos, dado que el consumo de memoria de las aplicaciones lingüísticas suelen ser considerables.

A la vez permite configurar cada instancia de su uso en forma independiente, mediante un conjunto de parámetros especializados enviados en cada requerimiento, pudiendo de este modo procesar múltiples idiomas y modos de operación en forma eficiente.

Búsqueda Inteligente en Bibliotecas



Se ideó un sistema de **Busqueda de Referencia Inteligente para Bibliotecas** y otras fuentes de cultura y conocimiento, permitiendo búsquedas inteligentes con capacidad conceptual, prescindiendo de ortografía estricta inclusive para nombres propios. Se usaron los algoritmos desarrollados: restauración fonética y ortográfica, lematización, tesauros ontológicos, a la vez de valoración por TF*IDF, etc. El acrónimo IOPAC proviene de "Internet Open Public Access Catalog"

Ej.: Análisis Morfológico Robusto

vimo hoi ezpe kavrom kon eza vavema ozpitalaria ke nempekano salu? klhdrdzkuic

[illegible]

Se muestra el resultado de un análisis morfológico de palabras muy mal escritas, las etiquetas son EAGLES 2.0, adicionadas con semántica, verosimilitud y probabilidad. La capacidad de restauración parece “casi humana”, hasta con 50-60% de letras mal.

Aplicaciones Potenciales

Asistencia/Enseñanza: Diálogo Robusto con los más Pequeños.
Reducir la Brecha Digital: Sistemas de Referencia Inteligentes.
Medicina: Historias Clínicas Médicas Digitales Interactivas.
Bases de Datos: Limpieza de Errores y Registros Duplicados.
Comprensión Artificial: Extracción/Inferencia Semántica.
Síntesis/Reconocimiento de Voz (usando diccionario morfológico).
Análisis Gramatical y Semántico Robusto (deep parsing).
P.L.N. Robusto, incluyendo OOV (Out Of Vocabulary words).

La creación y edición de un diccionario morfológico es una tarea compleja, pues requiere de ingresar las palabras, sus etiquetas, expresar sus reglas morfológicas en forma precisa, luego probar el resultado, importar y analizar listados de palabras por lotes, etc. Se han creado utilitarios muy específicos para poder crear, editar y enriquecer estos recursos con un mínimo esfuerzo.