

**Atribución de autoría: ¡El Móndrigo!
Bitácora del Consejo Nacional de Huelga**

Informe de Resultados

Elaborado para:
Gerardo Antonio Martínez, Confabulario

Autores:
Gerardo Sierra Martínez
Tonatiuh Hernández García
Brian Alfonso Sánchez Vázquez
Samuel Ramos Sosa

29 de abril de 2020

Contenido

Introducción	4
Contexto histórico del documento	4
Objetivo	7
Estructura del informe	7
Marco Conceptual	8
Atribución de autoría	8
Sistema Automático de Estudios Estilométricos (SAUTEE)	9
Características estilométricas	10
Distancias	14
Aprendizaje supervisado	17
Support Vector Machine	17
Logistic Regression	18
K-Neighbors	18
Decision Trees	18
Aprendizaje no supervisado	19
K-means:	19
Métricas de Distancia Textual	19
Método Kilgarriff	19
Método Delta de John Burrows	21
Resultados	22

Aprendizaje Supervisado	22
Clasificación Multiclase vía SAUTEE	22
Clasificación vía signos de puntuación	27
Aprendizaje no Supervisado	29
Clústers de características léxicas, y BOW	29
Métricas de distancia textual	30
Método Kilgarriff	30
Método Delta de John Burrows	31
Conclusiones	31
Referencias	32
Textos analizados	32
Bibliografía	34

Introducción

El Grupo de Ingeniería Lingüística atendió una solicitud para atribuir la autoría de un texto titulado: *¡El Móndrigo! Bitácora del Consejo Nacional de Huelga* de entre cinco autores probables.

Se cuenta con la transcripción del texto dubitado y cinco corpus de textos de autoría comprobada pertenecientes a los autores probables. Cuantificamos el estilo de escritura de todos los textos, con el fin de determinar estadísticamente si existe una similitud numérica que permita establecer una relación de autoría positiva entre *¡El Móndrigo!* con alguno de los autores candidatos.

Contexto histórico del documento

Poco después de la masacre de octubre de 1968 en la plaza de las Tres Culturas, la editorial fantasma Alba Roja publicó *¡El Móndrigo! Bitácora del Consejo Nacional de Huelga*.

El texto, de 182 páginas, fue presentado como la edición de un diario manuscrito perteneciente a un joven encontrado muerto al pie del edificio Chihuahua el 3 de Octubre y: “da cuenta de sucesos de máxima relevancia, puesto que había pertenecido a un líder del CNH” (Anónimo, 1968, pág. 5).

¡El Móndrigo! ha sido señalado como un libelo que dedica sus páginas a difamar y desinformar sobre el movimiento del 68, mezclando información real sobre asambleas y marchas con mentiras, burlas y difamaciones (Martré, 1998). Los libelos son “mecanismos empleados ya sea por el Estado para restarle fuerza a una figura política, para deslegitimar un movimiento social o para promover una campaña de manera páfida, entre otros objetivos no menos nefastos” (Belén, 2018, pág. 1). Publicado de manera anónima, no había ninguna

forma de saber quién había escrito el libelo, pues la editorial Alba Roja no existe (Tasso, 2018)

El 5 de octubre de 2003, en una entrevista, Sergio Romero Ramírez alias “El Fish”, porro, golpeador progubernamental, vinculado al MURO y al Yunque, aportó datos sobre la elaboración de libros que manipulaban y ocultaban la represión gubernamental:

“De la Secretaría de Gobernación salieron estos tres libros. Todos son de la misma editorial y son de la misma pluma, un hombre de paja, porque los autores no existen.

- ¿Quién escribió los libros entonces?

-La Secretaría de Gobernación, los repartía la Dirección Federal de Seguridad (DFS) y los responsables eran Fernando Gutiérrez Barrios, Miguel Nazar Haro y Luis de la Barreda.

-¿Pero específicamente quién se encargó de redactarlos?

-Jorge Joseph, que fue alcalde de Acapulco, por instrucciones de Gutiérrez Barrios.

Joseph fue alcalde de Acapulco a principios de los sesenta Muerto hace dos años, Joseph, dice El Fish, escribió esos libros con base en los informes de la DFS, con la que colaboraban los militares Arturo Acosta Chaparro y Humberto Quirós Hermosillo” (Delgado, 2003).

Según la investigación desarrollada por (Tasso, 2018) *¡El Mándrigo!* fue escrito por órdenes del capitán Fernando Gutiérrez Barrios, quién era director de la Dirección General

de Seguridad. El encargado de escribir el libelo era Jorge Joseph Piedra, un periodista veracruzano que formaba parte de un grupo de agentes de inteligencia que cubrían marchas, asambleas y actividades estudiantiles en Ciudad Universitaria y otras sedes. Piedra destacó con sus informes, lo cual lo hacía el candidato idóneo para encargarse de la escritura de *¡El Móndrigo!*

Según dicha investigación, Jorge Joseph utilizó los archivos policiales de los agentes a su cargo para dar verosimilitud al texto. Mezclándolos con su prosa, crea la ilusión de un diario meticuloso. Estos informes se pueden identificar en el texto por la mención de trayectorias de marchas, ubicación de asambleas, nombres de oradores, maestros y líderes del movimiento, cifras de heridos, detenidos, muertos y desaparecidos, detalles sobre la duración y temas de las asambleas, y cifras de asistentes a mítines y marchas (Tasso, 2018).

Otro autor es sospechoso de haber escrito el texto: Emilio Uranga. Gerardo Medina Valdés (1972), publicó que Emilio Uranga era el autor de *¡El Móndrigo!* A esta opinión se suma Gonzalo Martré (1998). Asimismo:

“En una carta que Octavio Paz escribió a Carlos Fuentes el 3 de agosto de 1969 señala a Emilio Uranga: ‘... se promueven y se publican ascos como El Móndrigo, un folleto obra de la cucaracha llamada [Emilio] Uranga’ (Sheridan, 2018).

A pesar de que ninguno de los autores mencionados presenta pruebas tangibles que permitan asociar a Uranga con la autoría de *¡El Móndrigo!*, lo cierto es que, en aquel entonces, el filósofo fungía como ideólogo para la Secretaría de Gobernación.

Las características de *¡El Móndrigo!*, que algunos autores como Munguía (2016) y Monsiváis (2018) han señalado como “diario apócrifo”, “mecanismo de estado para

deslegitimar el movimiento estudiantil” o “libelo difamatorio”, lo convierten en una obra controvertida y rodeada de polémica.

Las investigaciones actuales sobre este texto no han resuelto el problema de la autoría: algunas señalan al periodista veracruzano Jorge Joseph Piedra como el autor del texto (Redacción, 2003; Tasso, 2018). Otras fuentes señalan al filósofo Emilio Uranga, pero no aportan pruebas concretas sobre la relación de Uranga y el texto (Martré, 1998; Munguía, 2016; Sheridan, 2018). Los expertos no han llegado a un acuerdo final sobre la autoría del texto, que se mantiene anónimo.

Objetivo

Determinar la autoría de ¡El Mándrigo! utilizando métodos computacionales de atribución de autoría, lo que permite tener un enfoque objetivo e imparcial. Se utilizan los textos de autoría comprobada (indubitados) de cinco autores sospechosos entregados para crear su perfil estilométrico y compararlo con el estilo presente en el texto de autoría dubitada para determinar una relación de autoría basada en una metodología imparcial y científica.

Los cinco autores indubitados son: Emilio Uranga, Jorge Joseph Piedra, Roberto Blanco Moheno, Gregorio Ortega Molina y Gregorio Ortega Hernández. Los textos de cada uno de ellos se encuentran citados en la sección de referencias.

Estructura del informe

Este informe contiene una introducción en la que se detalla el contexto histórico del documento y se delimitan algunos autores sospechosos, coincidentes con los autores entregados para su evaluación.

Continuamos con el marco teórico que describe los principales conceptos requeridos para la comprensión técnica del informe. Describimos el funcionamiento teórico y práctico de los métodos aplicados.

La tercera parte de este informe está dedicada a la metodología. Teniendo en cuenta las diferencias de implementación de cada uno de los métodos presentados, nos centramos en los aspectos comunes: el preprocesamiento de los textos y la extracción de características.

En la cuarta parte, plasmamos los resultados de la aplicación de los distintos métodos. Todos los resultados se presentan de manera cuantitativa y se muestran gráficas cuando así lo permiten los datos.

Finalmente, en las conclusiones mostramos quién es el autor más probable con base en los resultados obtenidos por los diferentes métodos.

Marco Conceptual

A continuación, se introducen los conceptos necesarios para la comprensión de este informe técnico.

Atribución de autoría

En una atribución de autoría: “el autor de un texto puede ser seleccionado de un conjunto de posibles autores mediante la comparación de las medidas textuales del texto con las mismas medidas de la muestra de escritura del autor” (Grieve, 2007, pág. 1).

En 1994, en “Authorship attribution. Computers and the Humanities”, Holmes D.I. plantea que cuantificar el estilo literario utilizando distintas variables puede ser utilizado como una “huella digital” estilométrica. Para Holmes (1994, pág. 1), la estilometría se define como la búsqueda de unidades cuantificables que muestren con exactitud el estilo del texto,

que es un conjunto de patrones medibles que pueden ser únicos para cada autor. La atribución de autoría utiliza estas cuantificaciones con el fin de identificar y reconocer el estilo de un escritor.

La estilometría es la disciplina que cuantifica los textos. La atribución de autoría utiliza estos datos numéricos con el fin de identificar y reconocer el estilo de un escritor.

Holmes D.I. plantea que la cuantificación del estilo literario utilizando distintas variables puede ser utilizado como una “huella digital” estilométrica. Para Holmes (1994, pág. 1), la estilometría se define como: “la búsqueda de unidades cuantificables que muestren con exactitud el estilo del texto, que es un conjunto de patrones medibles que son únicos para cada autor”.

El estilo escrito es una combinación de decisiones consistentes en la producción del lenguaje en los diferentes niveles lingüísticos: elecciones léxicas, estructuras sintácticas, coherencia en el discurso (Daelemans, 2013). La estilometría es una disciplina que se encarga de cuantificar las características estilísticas de un texto que deben ser: “sobresalientes, estructurales, frecuentes, fácilmente cuantificables, y relativamente inmunes al control consiente del escritor” (Bailey, 1979, pág. 1). Estas características pueden ser utilizadas para identificar autores de textos anónimos o documentos en disputa (Stanczyk & A. Cyran, 2008).

Sistema Automático de Estudios Estilométricos (SAUTEE)

El Sistema Automático de Estudios Estilométricos (SAUTEE), es un sistema accesible en línea desde cualquier computadora con conexión a Internet, desarrollado en el Grupo de Ingeniería Lingüística, que permite al usuario identificar las similitudes en un conjunto de documentos con base en distintos marcadores estilométricos (López-Escobedo, Sierra, &

Solórzano, 2019). Para ello, el sistema proporciona hasta 26 características estilométricas que incluyen n-gramas de caracteres, de palabras, de categorías gramaticales, uso de signos de puntuación y riqueza de vocabulario.

Características estilométricas

1) Bigramas de caracteres

Se contabilizan las apariciones de dos caracteres seguidos. Los espacios se consideran caracteres. La frecuencia de cada bigrama se divide entre el total de apariciones contabilizadas en el texto.

2) Trigramas de caracteres

Lo mismo que bigramas pero considerando tres caracteres seguidos.

3) Unigramas de palabras funcionales

Se contabiliza la aparición de las palabras funcionales (artículos, determinantes, preposiciones, conjunciones, nexos). La frecuencia de cada palabra se divide entre el total de apariciones.

4) Bigramas de palabras funcionales

Contabiliza todas aquellas apariciones de dos palabras funcionales seguidas. La frecuencia de cada bigrama se divide entre el total de apariciones.

5) Bigramas de palabras funcionales con hasta 2 huecos

En este caso se contabilizan las apariciones de dos palabras funcionales que no están contiguas, sino separadas a lo más, por otras dos palabras.

6) Trigramas de palabras funcionales

Los mismo que bigramas de palabras funcionales, pero tomando en cuenta las apariciones de tres palabras funcionales seguidas.

7) Trigramas de palabras funcionales con hasta 2 huecos

Lo mismo que bigramas de palabras funcionales con dos huecos, pero considerando apariciones de tres palabras funcionales. No necesariamente tiene que haber el mismo número de huecos entre la primera y la segunda palabra funcional que entre la segunda y la tercera.

8) Longitud de oraciones

Se contabiliza las frecuencias de las siguientes categorías de oraciones: menos de 10 palabras, 11 a 20, 21 a 30, 31 a 40, 41 a 50 y más de 50. Cada frecuencia se divide entre el número total de oraciones en el texto.

9) Longitud de palabras

Se contabiliza la aparición de las siguientes categorías de palabras, de 1 letra, de 2 letras y hasta llegar a 20 letras. Cada frecuencia se divide entre el número total de palabras en el texto.

10) Unigramas de etiquetas POS

Se contabiliza la aparición de etiquetas “Parts of Speech” (partes del discurso) tal como Freeling (Padró, 2011) las genera. Freeling es una librería en C++ que proporciona funcionalidades para el análisis de lenguaje. Fue desarrollado por TALP Research Center en la Universitat Politècnica de Catalunya. La frecuencia de cada etiqueta se divide entre el número de palabras en el texto.

11) Bigramas de etiquetas POS

Se contabilizan las apariciones de dos etiquetas POS contiguas. La frecuencia de cada bigrama se divide entre el número de palabras en el texto.

12) Trigramas de etiquetas POS

Se contabilizan las apariciones de tres etiquetas POS contiguas. La frecuencia de cada trigramas se divide entre el número de palabras en el texto.

13) Categoría gramatical al inicio de la oración

A partir de las etiquetas POS generadas por Freeling, se contabiliza el número de veces que cada categoría gramatical aparece al inicio de la oración. Cada frecuencia es dividida entre el número total de palabras al inicio de la oración.

14) Categoría gramatical al final de la oración

A partir de las etiquetas POS generadas por Freeling, se contabiliza el número de veces que cada categoría gramatical aparece al final de la oración. Cada frecuencia es dividida entre el número total de palabras al final de la oración.

15) Unigramas de etiquetas POS no fino

Lo mismo que unigramas de etiquetas POS, pero en vez de contabilizar la frecuencia de las etiquetas tal cual las genera Freeling, se toma en cuenta únicamente el primer carácter de la etiqueta, que corresponde a la categoría gramatical más general.

16) Razón de adjetivos

Frecuencia de adjetivos por cada 100 palabras.

17) Razón de sustantivos

Frecuencia de sustantivos por cada 100 palabras.

18) Razón de preposiciones

Frecuencia de preposiciones por cada 100 palabras.

19) Razón de verbos

Frecuencia de verbos por cada 100 palabras.

20) Bigramas de etiquetas POS no fino

Lo mismo que unigramas de etiquetas POS no fino, pero contabilizando las apariciones de dos etiquetas seguidas. La frecuencia de cada bigrama se divide entre el total de apariciones contabilizadas en el texto.

21) Trigramas de etiquetas POS no fino

Lo mismo que bigramas de etiquetas POS no fino, pero contabilizando las apariciones de tres etiquetas seguidas. La frecuencia de cada trigramas se divide entre el total de apariciones contabilizadas en el texto.

22) Signos de puntuación

Se contabilizan los signos de puntuación del texto con base en el etiquetado Freeling (es decir, se toma como signo de puntuación todo lo que Freeling etiqueta como tal).

Cada frecuencia se divide entre el número total de signos de puntuación en el texto.

23) Riqueza de vocabulario - W de Brunét

Descripción: W de Brunét

$$W = N^{V^{-0.165}}$$

Donde N = número total de palabras,
 V = Tamaño del vocabulario

24) Riqueza de vocabulario - R de Honoré

Descripción: R de Honoré

$$R = \frac{100 \times \log N}{1 - \frac{V_1}{V}}$$

Donde
 N = número total de palabras
 V_1 = número de palabras que aparecen una sola vez
 V = tamaño del vocabulario

25) Riqueza de vocabulario - Medida de Sichel

Descripción: Medida de Sichel

$$S = \frac{V_2}{V}$$

Donde

V_2 = número de palabras que aparecen dos veces

V = tamaño del vocabulario

Distancias

SAUTEE cuenta las apariciones de las características seleccionadas y normaliza las frecuencias. Para poder determinar qué tan parecidos son cada par de documentos entre sí, crea vectores por cada documento y determina su semejanza mediante una medida de distancia (López Escobedo, Solorzano Soto, & Sierra, 2016). Entre las diferentes fórmulas para medir la distancia se encuentran, en SAUTEE se manejan cuatro: Euclidiana, Delta, Manhattan y Canberra.

La distancia euclidiana es igual a la raíz cuadrada de la suma del cuadrado de las diferencias de cada dimensión:

$$\sqrt{\sum_{i=1}^n (X_i - Y_i)^2}$$

La Delta, tal y como fue definida por Burrows (2002), es el promedio de las diferencias absolutas entre los z-scores de un conjunto de palabras en un grupo de textos y los z-scores del mismo conjunto de palabras en el texto objetivo.

$$\sum_{i=1}^n \left| \frac{X_i - Y_i}{\sigma_i} \right|$$

La distancia Manhattan es también llamada distancia de taxista, haciendo referencia a la distancia que un vehículo tendría que recorrer para llegar de un punto a otro en una cuadrícula. Matemáticamente es igual a la suma de las diferencias absolutas entre cada dimensión.

$$\sum_{i=1}^n |X_i - Y_i|$$

La distancia Canberra una distancia Manhattan ponderada. La diferencia absoluta entre las variables es dividida entre la suma de los valores absolutos de las mismas. Es igual a:

$$\sum_{i=1}^n \frac{|X_i - Y_i|}{|X_i| + |Y_i|}$$

Esta distancia tiene la bondad de ser sensible a valores cercanos a cero, por lo cual es útil si el conjunto de datos contiene tanto valores pequeños como valores muy grandes (López Escobedo, Solorzano Soto, & Sierra, 2016).

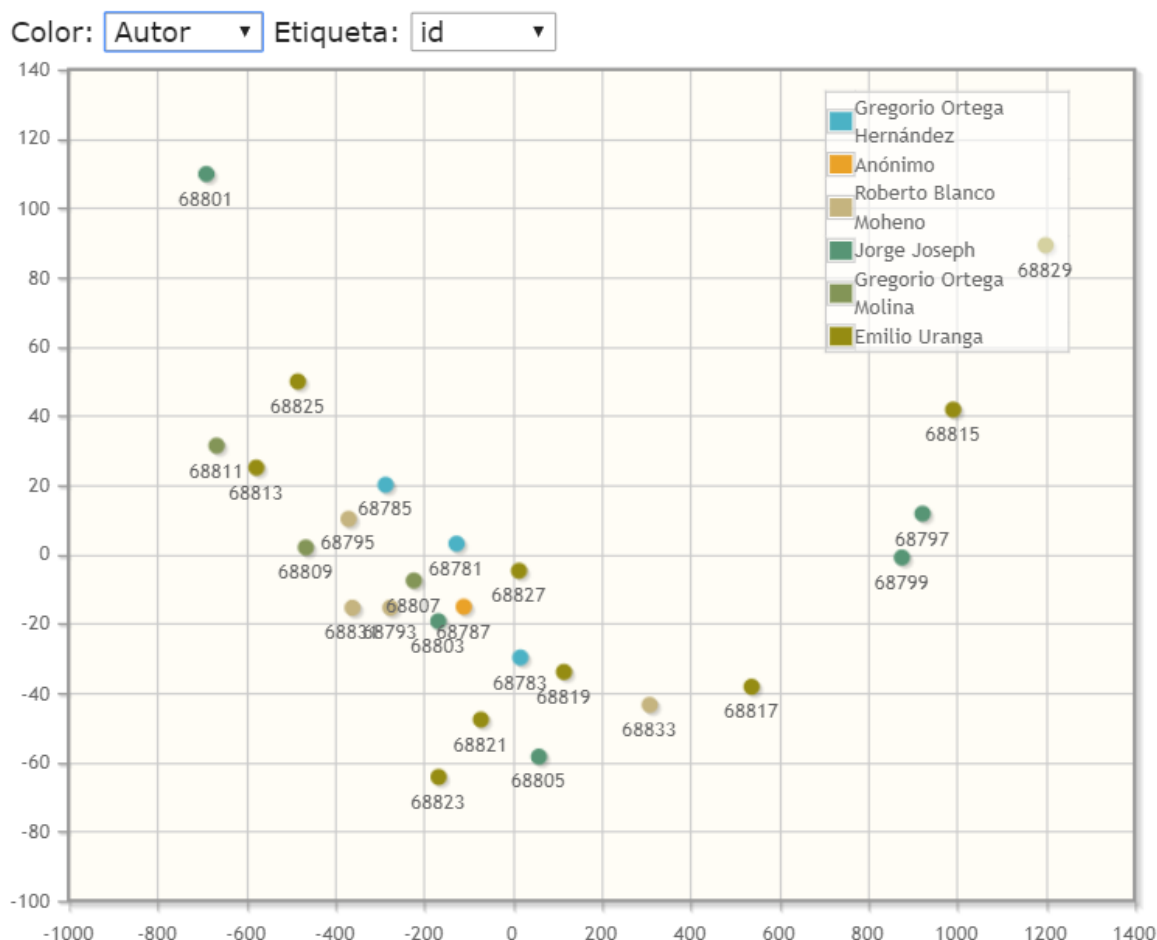
El sistema SAUTEE está diseñado para que, a partir de la extracción de marcadores estilométricos, dado un conjunto de documentos, vectorice los textos y con ello calcule la distancia entre pares de documentos, para finalmente “generar un conjunto de puntos en dos

dimensiones, por medio del cual se pueden visualizar en una gráfica las distancias calculadas” (López-Escobedo, Sierra, & Solórzano, 2019).

Por sí mismo, SAUTEE es capaz de identificar textos cercanos entre sí de acuerdo con la cercanía entre los vectores estilométricos que los representan. En la gráfica 1, se observan los textos identificados por un id y el color del autor que los ha escrito.

El Móndrigo está identificado por el id 68787, con el color naranja. Es visible la cercanía con el punto identificado con el id 68803, que corresponde a Jorge Joseph. Esto significa que los vectores de características de estos textos están muy cercanos, lo que puede señalar una relación de autoría positiva entre el Móndrigo y Jorge Joseph.

Gráfica 1



Sin embargo, para obtener una certeza matemática se utilizarán los datos obtenidos mediante la extracción de características de SAUTEE para implementar algoritmos que nos permitan identificar con mayor precisión el autor del Mándrigo.

Aprendizaje supervisado

El Aprendizaje de Máquina es una disciplina que se enmarca en la Inteligencia Artificial que crea “sistemas que aprenden automáticamente” (González, 2019). En este contexto, “aprender significa identificar patrones en los datos”, “la máquina es un algoritmo matemático que revisa los datos y es capaz de predecir comportamientos futuros” (González, 2019) o en nuestro caso, indicarnos la similitud entre dos patrones de datos que son las ocurrencias de las características estilométricas seleccionadas. El aprendizaje supervisado utiliza “etiquetas” en cada instancia con que se entrena el sistema. En este caso, los datos numéricos que caracterizan a cada autor tienen una etiqueta con el nombre del autor para que el sistema asocie esos datos con el nombre de quien le corresponde.

En una atribución de autoría donde un texto de autoría desconocida es asignado a un autor de un set de autores candidatos con textos de autoría probada, desde el punto de vista del aprendizaje supervisado de máquina, es visto como una tarea de clasificación multiclase, de etiqueta única y de clasificación de texto (Sebastiani, 2002).

Support Vector Machine

Para la máquina de soporte vectorial (SVM), el conjunto de datos de entrenamiento etiquetado sirve para construir un modelo que predice la posibilidad de que una nueva clase pertenezca a una categoría definida previamente por el entrenamiento. SVM construye con los datos de los autores una nube de puntos en conjuntos de hiperplanos y busca la mejor separación entre las clases etiquetadas. Al introducir datos nuevos (los datos de ¡El

Móndrigo!) SVM es capaz de clasificarlos en las categorías entrenadas. Cuando SVM recibe los datos estilométricos, es capaz de predecir con qué probabilidad pertenece a uno de los cinco autores.

Logistic Regression

Es un tipo de análisis que calcula las relaciones entre variables. Se utiliza para predecir el valor de una variable en función de las variables independientes, modelando la probabilidad de que ocurra un evento considerando otros factores. Entrenado con las medidas estilométricas de los autores, estimará la probabilidad de que los nuevos datos pertenecientes al texto anónimo pertenezcan a un autor.

K-Neighbors

Es un método de clasificación supervisada. Dado un conjunto de datos de entrenamiento, se entrena un clasificador no paramétrico que estima la probabilidad de que los nuevos datos del texto anónimo pertenezcan a una clase (los autores candidatos). Los datos de entrenamiento son vectores multidimensionales, cada párrafo tiene atributos numéricos asignados a una clase. El espacio es dividido en regiones y cuando se recibe un dato nuevo para clasificarlo se calcula la distancia entre los vectores, se seleccionan los ejemplos más cercanos y el ejemplo nuevo es asignado a la clase que más coincidencias tiene.

Decision Trees

Dado un conjunto de datos de entrenamiento, se forma un modelo predictivo basado en diagramas de construcciones lógicas que categorizan condiciones de forma sucesiva: al recibir los datos del texto anónimo, comprobará qué elecciones estilométricas (representadas en números) coinciden con el estilo de los escritores candidatos, asignando cada párrafo al

autor más probable, tomando en cuenta todas las “elecciones” contenidas (comas, longitud de párrafo, palabras por párrafo etc.).

Aprendizaje no supervisado

K-means:

Es un método de agrupamiento que obtiene la partición de un conjunto de n observaciones en k grupos en el que cada observación pertenece al grupo cuyo valor medio es más cercano. K-medias encuentra valores de características semejantes y las coloca juntas. Dado un conjunto de características propias de un documento X , al presentarse un grupo de documentos, el algoritmo los “colocará” cercanos en una gráfica.

Métricas de Distancia Textual

Las métricas de distancia textual fueron los primeros métodos utilizados para cuantificar los textos y determinar su semejanza con base en patrones numéricos. A continuación se detalla su funcionamiento.

Método Kilgariff

Para resolver el problema de la autoría se utilizó el método Kilgariff, que establece el concepto de *distancia* entre textos: si un autor escribe dos textos y ambos se comparan, la distancia entre ellos debe ser mínima, lo que significa que ambos textos han sido escritos por la misma persona (Kilgariff, 2001).

Kilgariff hace una comparación entre distintos tipos de medidas para semejanza de corpus utilizando conteo de palabras, y lleva a cabo sus experimentos dividiendo el texto de un autor en secciones y mezclándolas en dos corpus, para medir las diferencias estadísticamente. Su conclusión es que la mejor medida es la llamada X^2 (fórmula 1) que es

un índice típicamente utilizado para medir independencia entre variables aleatorias categóricas (es decir, no numéricas).

$$\chi^2 = \sum_i \frac{(C_i - E_i)^2}{E_i} \quad (1)$$

El detalle es que en los experimentos de Kilgariff se ocupa solo un texto dividido, y para lograr un efecto similar lo que procede es que se toman los n términos más comunes de cada corpus y se calcula el valor esperado de las ocurrencias de cada una de estas n palabras como si los dos corpus fueran muestras aleatorias de la misma población.

Si los tamaños de los corpus 1 y 2 fueran respectivamente N_1 Y N_2 , y la palabra w ha sido observada respectivamente con frecuencias o_{w1} y o_{w2} , entonces los valores esperados de ocurrencia de la palabra w en cada uno de los corpus serán los mostrados en las siguientes ecuaciones:

$$e_{w1} = \frac{N_1 \times (o_{w1} + o_{w2})}{N_1 + N_2}$$

$$e_{w2} = \frac{N_2 \times (o_{w1} + o_{w2})}{N_1 + N_2}$$

Para su implementación, se agrupan primero todos los textos del autor en un solo texto y se convierte ese texto en una lista de términos que incluyen palabras y puntuación. Entonces, se elimina la puntuación y para cada autor se reúne su corpus y el corpus del texto anónimo en una sola lista.

Se encuentra la ocurrencia de cada término, se eligen los n más frecuentes y se calcula el valor esperado de cada palabra de las n más frecuentes tomando la contribución de cada autor a este corpus conjunto. Ello se calcula con el número de términos del corpus de cada uno

entre el número de términos del corpus conjunto y multiplicándolo por el número de ocurrencias de esa palabra en el corpus de cada autor.

Se calcula para cada palabra un término cuyo valor es el cuadrado de la diferencia entre ese valor esperado y el número de ocurrencias en el corpus del autor, dividido entre el valor esperado. Finalmente, X^2 es la suma de todos estos términos, de acuerdo con la fórmula 1.

Método Delta de John Burrows

Sofisticando el concepto de distancia, Burrows (2002) propone un método que, dado un documento de prueba y una serie de documentos cuya autoría está plenamente definida, arroja una medida de distancia con todos los autores, identificando un estilo particular de cada autor y midiendo la distancia entre ese estilo y el texto en disputa, el de prueba. Esta medida se llama Δ (delta) y, entre más pequeña, más certeza habrá del origen de ese texto.

Para implementar esta medida, en primer lugar se agrupan todos los textos en un corpus. Se buscan las palabras más frecuentes para usarlas como atributos o características (dimensiones). Se encuentra la proporción con la que contribuye cada autor a esa característica. Se calcula el promedio y la desviación estándar para usarlas como atributos del corpus conjunto. Para cada uno de los atributos (palabras) y usando los valores calculados en el paso anterior, se calcula el valor Z donde C_i es el promedio de los promedios como la frecuencia de esa palabra en el corpus y aplicando la siguiente fórmula (2):

$$Z_i = \frac{C_i - \mu_i}{\sigma_i} \quad (2)$$

Para calcular la Δ se compara el texto de cada autor con el corpus en disputa de acuerdo a fórmula 3, donde los subíndices $c(i)$ y $t(i)$ se refieren al atributo del autor y de prueba, respectivamente.

$$\Delta_c = \sum_i \frac{|Z_c(i) - Z_t(i)|}{n} \quad (3)$$

Resultados

A continuación se presentan los resultados de los diferentes métodos implementados.

Aprendizaje Supervisado

Clasificación Multiclase vía SAUTEE

En el método Clasificación multiclase OVR en vectores de características obtenidas con SAUTEE, se abordó la atribución de autoría del texto anónimo como un problema de clasificación de texto, en donde las características estilométricas son cuantificadas y utilizadas como clases. Los valores fueron utilizados para entrenar una Support Vector Machine que predice los valores para cada clase (Koppel , Schler, & Argamon, 2009).

Se realizó el entrenamiento con los vectores generados por la cuantificación estilométrica del SAUTEE. Las categorías medidas se encuentran en la tabla 1. Se utilizó un algoritmo SVC (Support Vector Classifier) implementado en LIBSVM de MatLab con el paradigma One vs Rest, una clasificación binaria que hace pares con todos los elementos del conjunto para minimizar errores de clasificación y obtener los mejores parámetros para la generación del modelo. Se verificó la precisión del entrenamiento sobre un corpus segmentado de los

mismos autores. El modelo obtenido se compara con los valores del texto dubitado, obteniendo un porcentaje de similitud entre la SVM entrenada con los estilos de escritura de los autores y ¡El Mándrigo! indicando la probabilidad de que un autor hubiese escrito el texto (Chung & Chih-Jen , 2011).

Tabla 1: Características de SAUTEE medidas para su análisis en SVM

ID_MARCADOR	MARCADOR
BC	Bigramas de caracteres
TC	Trigramas de caracteres
UPF	Unigramas de palabras funcionales
BPF	Bigramas de palabras funcionales
BPF2H	Bigramas de palabras funcionales con hasta dos huecos
TPF	Trigramas de palabras funcionales
TPF2H	Trigramas de palabras funcionales con hasta dos huecos
LO	Longitud de oraciones
LP	Longitud de palabras
UPOS	Unigramas de etiquetas POS
BPOS	Bigramas de etiquetas POS
TPOS	Trigramas de etiquetas POS
CGIO	Categoría gramatical al inicio de la oración
CGFO	Categoría gramatical al final de la oración
UPOSNF	Unigramas de etiquetas POS no fino

Se seleccionaron aleatoriamente doce textos de veinticuatro pertenecientes a los autores probables para realizar el entrenamiento. El resto de los textos fue utilizado como datos de

prueba o validación. Se usó validación cruzada para obtener los mejores parámetros de kernel para el clasificador y generar el modelo óptimo para estimar la probabilidad de similitud entre cada clase de los vectores que representan numéricamente al texto dubitado. En la tabla 2 se muestran los parámetros buscados con este método.

Tabla n. 2: Parámetros del algoritmo SVM

K	Función de Kernel, se contemplan cuatro: Lineal, Polinomial, RBF, y Sigmoide.
C	Es la función de costo: a menor costo, se tiene cierta holgura para hallar un hiperplano de mejor ajuste.
γ	Es el coeficiente de la función del Kernel.
D	Degree o grado de la función del Kernel, para el Kernel polinomial, irrelevante para las demás.

Se corrieron 100 iteraciones por función de Kernel y marcador estilométrico. Se contemplaron 4 funciones de Kernel: lineal, polinomial, rbf y sigmoideal, se contemplaron también para hallar el grado de ajuste de la función de kernel polinomial del grado 3 al 7.

En el caso que hubiese empate, se consideró el kernel que tuviera una función de costo menor, esto para evitar la penalización y tener el mayor número de vectores de soporte. Una vez obtenidos los mejores parámetros para el modelo del clasificador, se corrieron dichos parámetros del kernel sobre los vectores obtenidos de *¡El Mándrigo!* La precisión de los parámetros se pueden consultar en la tabla 3.

Se seleccionan aquellas funciones y parámetros kernel de mayor de precisión, para cada marcador de estilométrico y en caso de las funciones de kernel empatadas, se selecciona aquella con una función de penalización C para que el clasificador tenga una mayor holgura al buscar el mejor hiperplano de separación.

Tabla 3: Parámetros obtenidos con distintas funciones de Kernel

ID_MARC	LINEAL	POLI	RBF	SIGMOIDE
BC	83.3333%: C=8.724, gamma=0.15328 d=3	83.3333%: C=32, gamma=0.014148, d=3	83.3333%: C=279.17, gamma=0.02181 9	83.3333%: C=181.02, gamma=0.19035
TC	83.3333%:C=76.109, gamma=0.00783 1	75%: C=5.6569, gamma=0.0001586 8	91.6667%:C=76.109, gamma=0.02872 4	83.3333%:C=181.019, gamma=0.00138 43
UPF	75%: C=32, gamma=0.01098 9	83.3333%:, C=32, gamma=0.026136, d=3	75%: C=279.17, gamma=0.00034 341 d=4	75%: C=430.539, gamma=0.00014 438
BPF	75%: C=76.1093, gamma=0.01110 1	83.3333%:C=76.109, gamma=0.0019624	75%: C=430.539, gamma=0.00053 5	83.3333%:C=13.4543, gamma=0.00082 508
BPF2H	58.3333%:C=5.6569, gamma=0.07891	66.6667%:C=5.656, gamma=0.013949 d=3	75%: C=181.0193, gamma=0.01394 9	83.3333%:C=76.1093, gamma=0.00246 59
TPF	66.6667%: C=5.6569, gamma=0.00804 69	66.6667%: C=13.4543, gamma=0.0052177 d=3	66.6667%: C=76.1093, gamma=0.00338 33	66.6667%: C=181.0193, gamma=0.00338 33
TPF2H	66.6667%: C=13.4543, gamma=0.01857	66.6667%: C=13.454, gamma=0.0052177 d=3	66.6667%: C=76.109, gamma=0.00338 33	66.6667%: C=181.02, gamma=0.00338 33
LO	66.6667%: C=1, gamma=0.00520 83	75%: C=20.7494, gamma=0.0021898, d=3	75%: C=32, gamma=0.00520 83	75%: C=13.4543, gamma=0.00092 071
LP	75%: C=5.6569, gamma=0.05	83.3333%: C=49.350, gamma=0.021022 d=4	83.3333%: C=20.749, gamma=0.01363 1	83.3333%: C=76.109, gamma=0.02102 2
UPOS	83.3333%: C=8.7241, gamma=0.02459 5	83.3333%: C=32, gamma=0.024595, d=4	83.3333%: C=49.3507, gamma=0.00832 7	75%: C=0.013139, gamma=0.01034 1,
BPOS	83.3333%: C=32, gamma=0.00124 52	91.6667%: C=13.4543, gamma=0.0002201 2, d=4	75%: C=32, gamma=0.00124 52	75%: C=1, gamma=0.00022 012
TPOS	75%: C=181.0193, gamma=0.04052 9	75%: C=181.0012, gamma=0.0071645, d=3	58.3333%: C=32, gamma=0.00195 32	58.3333%:C=117.376, gamma=0.00126 65
CGIO	75%: C=117.3765, gamma=3.3116	83.3333%:C=13.454, gamma=1.7291, d=3	83.3333%: C=32, gamma=0.08333 3	83.3333%: C=2.3784, gamma=6.3424
CGFO	75%: C=2.3784, gamma=0.1, d=3	75%: C=20.7494, gamma=3.2, d=4	75%: C=13.4543, gamma=0.23784	75%: C=5.6569, gamma=3.2
UPOSNF	83.3333%: C=76.1093, gamma=0.00284 09	83.3333%: C=181.0193, gamma=0.0005022 1, d=3	83.3333%: C=2.3784, gamma=0.01607 1	75%: C=5.6569, gamma=0.09090

Se hizo un ranking con los mejores resultados de las funciones de kernel, identificando en el entrenamiento el que obtuviera la mayor precisión (tabla 4).

Tabla 4: Predicción de autoría vía SVM y marcadores estilométricos

ME	K	ORTEGA MOLINA	EMILIO URANGA	BLANCO MOHENO	ORTEGA HERNANDEZ	JORGE JOSEPH
BC	LIN	10.61%	3.94	1.82	14.01	69.62
TC	RBF	8.28	9.43	4.25	11.72	66.31
UPF	POLY	10.05	8.85	6.90	13.30	60.91
BPF	SIG	16.74	17.45	12.10	14.24	39.48
BPF2H	SIG	9.67	24.36	14.59	9.19	42.19
TPF	LIN	9.02	29.94	11.41	10.36	39.26
TPF2H	RBF	9.16	35.30	15.41	17.47	22.66
LO	SIG	14.70	42.58	8.80	12.37	21.55
LP	RBF	9.40	18.15	3.58	14.15	54.72
UPOS	LIN	15.36	3.43	7.54	14.02	59.65
BPOS	POLY	19.59	13.22	13.96	16.44	36.79
TPOS	POLY	32.08	8.08	14.42	29.59	15.83
CGIO	SIG	6.73	44.50	9.35	19.41	20.01
CGFO	LIN	13.10	36.80	5.99	17.28	26.83
UPOSNF	RBF	13.04	12.51	15.63	12.12	46.71

Dadas las características que se presentan en las tablas de resultados, se puede concluir que es muy posible que la limitada cantidad de datos esté provocando un sesgo, ya que la máxima precisión para la mayoría de los modelos es del 83.3%, a excepción del caso en

donde se utilizó como vectores de características aquellos generados con el marcador trigramas de caracteres y un kernel RBF, dando una confiabilidad del 91.6%.

Con base en la tabla de resultados, es altamente probable que ¡El Mándrigo! haya sido escrito por Jorge Joseph, pues coincide en la mayoría de las categorías medidas con porcentajes elevados, lo que indica una similitud muy elevada entre los estilos de escritura.

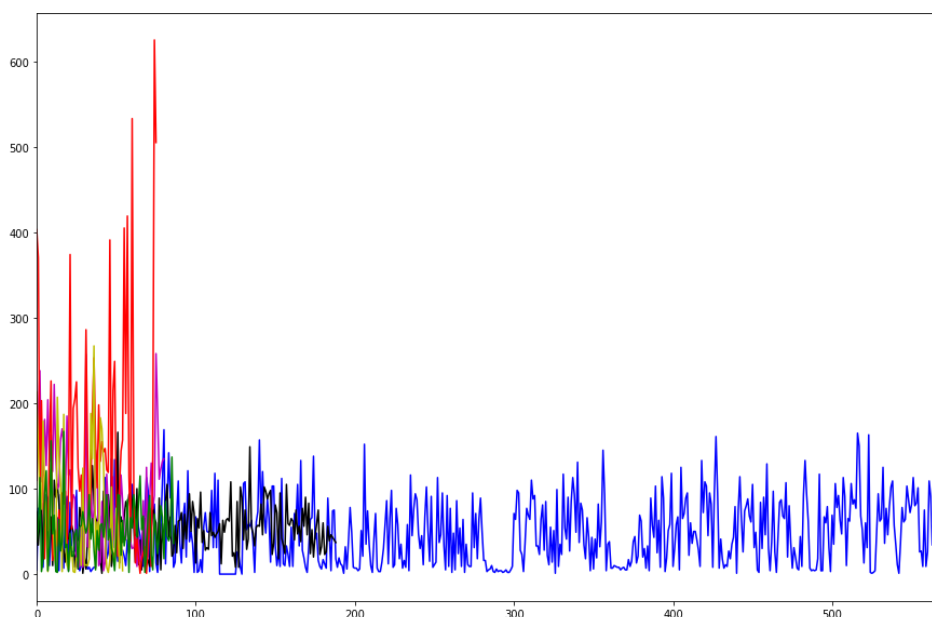
Clasificación vía signos de puntuación

Se utilizaron únicamente signos de puntuación para caracterizar el perfil de cada autor. El algoritmo indica la cantidad de párrafos que pudieron haber sido escrito por el autor indicado.

Este método, al emplear un número reducido de características, permite realizar visualizaciones de los datos. A continuación se muestran las gráficas que muestran el uso de signos de puntuación por autor: En azul se aprecia el Mándrigo, en negro a Jorge Joseph, en magenta a Emilio Uranga, a Blanco Moheno en rojo, Ortega en amarillo y Molina en verde.

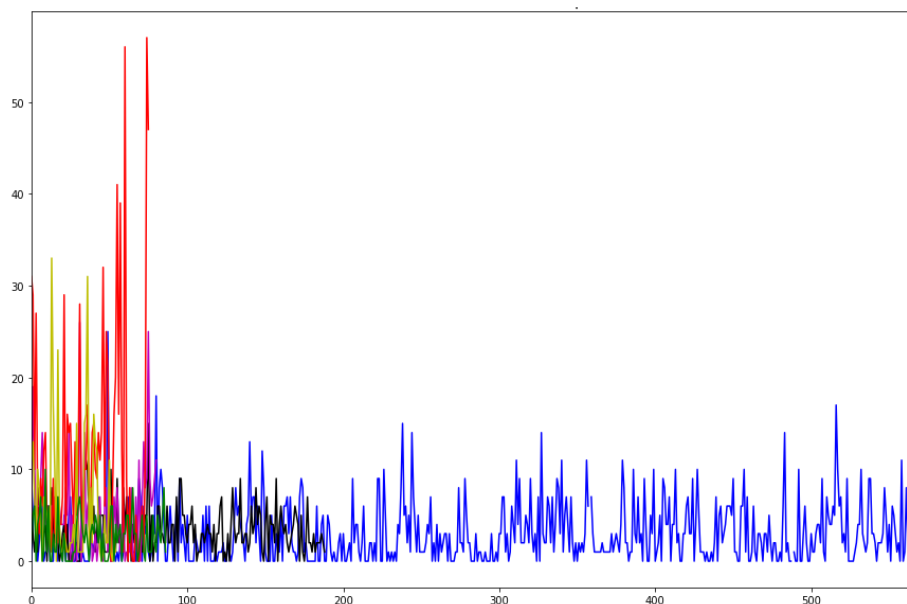
La gráfica 2 representa el uso del número de palabras en los corpus. Se puede observar que el color negro y el azul son muy similares, corresponden al Mándrigo y a Jorge Joseph.

Gráfica 2: Uso de palabras por autor



La gráfica 3 muestra el uso de comas:

Gráfica 3: uso de comas por autor



Nuevamente es consistente que las líneas negras y azules sean similares. La longitud de las líneas no es similar, pues el texto del Mórdrigo es más extenso que cualquier corpus de los autores, sin embargo, aunque esto es evidente en la gráfica, no representa ninguna dificultad para los algoritmos.

A continuación, en la tabla 5 se muestran los resultados de los algoritmos que clasificaron el número de párrafos del Mórdrigo que son atribuido a cada autor según el estilo de escritura presentes en ellos.

Tabla 5: Clasificación multiautor con signos de puntuación

Algoritmo	Precisión	Resultado
Logistic Regression	0.6010928961748634	Counter({' JorgeJoseph ': 330, 'BlancoMoheno': 187, 'OrtegaMolina': 23, 'EmilioUranga': 22, 'GregorioOrtega': 7})
Support Vector Machine SVC	0.7950819672131147	Counter({' JorgeJoseph ': 491, 'BlancoMoheno': 40, 'EmilioUranga': 18, 'OrtegaMolina': 16, 'GregorioOrtega': 4})
KNeighbors	0.6229508196721312	Counter({' JorgeJoseph ': 315, 'BlancoMoheno': 85, 'EmilioUranga': 77, 'OrtegaMolina': 67, 'GregorioOrtega': 25})
DecisionTree	0.9754098360655737	Counter({' JorgeJoseph ': 232, 'GregorioOrtega': 101, 'EmilioUranga': 87, 'OrtegaMolina': 79, 'BlancoMoheno': 70})

Todos señalan a Jorge Joseph como autor del texto.

Aprendizaje no Supervisado

Clústers de características léxicas, y BOW

Los clústers vía *K-means* agruparon en la misma categoría los vectores BOW y Léxico extraídos de los textos de ¡El Mándrigo! y Jorge Joseph, lo que indica una relación de autoría probable, basada en coincidencias léxicas (Tabla 6).

Tabla 6: Clústers de características BOW y Léxicas

	Uranga	OrtegaM	JorgeJoseph	GregorioO	RobertoB	Móndrigo
Léxico	Clúster1	Clúster2	Clúster0	Clúster4	Clúster5	Clúster0
BOW	Clúster1	Clúster2	Clúster0	Clúster4	Clúster5	Clúster0

Métricas de distancia textual

Método Kilgariff

El resultado final de comparar los textos de cada autor con el texto anónimo de acuerdo al método Kilgariff tomando las 500 palabras más comunes se muestran en la tabla 7 donde se observa que de acuerdo a la menor distancia (X^2) es más probable que el autor del texto anónimo haya sido Jorge Joseph, seguido de cerca por Emilio Uranga, en tanto el menos probable es Roberto Blanco Moheno.

Tabla 7: Método Kilgariff

Autor	Distancia X^2
Emilio Uranga	1771.0364889378193
Gregorio Ortega Molina	2002.4749684727008
Jorge Joseph	1643.598487830984
Gregorio Ortega Hernández	1986.895626691162
Roberto Blanco Moheno	3221.980751878392

Método Delta de John Burrows

El resultado final de comparar los textos de cada autor con el texto en disputa de acuerdo con método DELTA de Burrows es el que se muestra en la tabla 8. Podemos observar que el autor que tiene el menor DELTA y por tanto el que mayor con probabilidad es el autor del texto es Jorge Joseph.

Tabla 8: Delta de Burrows

Autor	Delta
Emilio Uranga	1.5134148360747048
Gregorio Ortega Molina	1.0422030585665234
Jorge Joseph	0.5190645312503361
Gregorio Ortega Hernández	1.2069074912312774
Roberto Blanco Moheno	1.2278603090115858

El autor con el valor Δ menor (más parecido al Mándrigo) es Jorge Joseph. Con base en esta medida, Joseph es el autor de ¡El Mándrigo!

Conclusiones

Se realizaron cinco análisis independientes. Todos los análisis implementados son coincidentes en señalar a Jorge Joseph Piedra como el autor con mayor similitud al estilo de escritura presente en ¡El Mándrigo!

La evidencia numérica indica una relación de autoría positiva entre Jorge Joseph Piedra y ¡El Mándrigo! Bitácora del Consejo Nacional de Huelga.

No hay evidencia suficiente o consistente para señalar a ningún otro autor.

Referencias

Textos analizados

El Móndrigo

Anónimo. (1968). ¡El Móndrigo! Bitácora del Consejo Nacional de Huelga. México: Alba Roja.

Gregorio Ortega Hernández

Hernández, G. O. (1968). México resisitíó a la CIA y a Fidel Castro. Revista de América.

Hernández, G. O. (1968). Un chamaco de Tepito da una lección de civismo. Revista de América.

Hernández, G. O. (1968). Un injusto artículo de Blanco Moheno. Revista de América.

Roberto Blanco Moheno

Moheno, R. B. (1966). La noticia detrás de la noticia. 180-190.

Moheno, R. B. (1966). La noticia detras de las noticias . 180-190.

Moheno, R. B. (1966). Tlatelolco. Historia de una infamia . La noticia detrás de la noticia , 103-105.

Moheno, R. B. (1995). La noticia detrás de la noticias. Revista de Noticias, 103-105.

Gregorio Ortega Molina

Molina, G. O. (1968). ¿Hacia la extrema derecha? . Revista de América.

Molina, G. O. (1968). La guerra de los zombies . Revista de América.

Molina, G. O. (1968). La sociedad del átomo . Revista de América.

Jorge Joseph Piedra

Piedra, J. J. (1961). Aviso. En J. J. Piedra, El ministro del odio (pág. 125). México: Edición del autor.

Piedra, J. J. (1961). Picaluga vende a los Rojos Gomistas por una Curul Senatorial . En J. J. Piedra, El ministro del odio (págs. 35-40). México: Edición del autor.

Piedra, J. J. (1961). Pitoloco en Palacio . En J. J. Piedra, El ministro del odio (págs. 119-124). México: Edición del autor .

Piedra, J. J. (1965). Esoterismo en Anáhuac. En J. J. Piedra, México, cuna de la civilización universal (págs. 263-277). México : Ramírez Editores.

Piedra, J. J. (1965). Pruebas de la existencia de la Atlántida. En J. J. Piedra, México: cuna de la civilización universal (págs. 87-96). México: Ramírez Editores .

Emilio Uranga

Uranga, E. (¿? de ¿? de 1968). La Ambigüedad Universitaria. La Prensa, pág. ¿?

Uranga, E. (¿? de ¿? de 1968). La Represión Extremada . La Prensa , pág. 3 y 38.

Uranga, E. (¿? de ¿? de 1968). Ornato y Orden . La Prensa, pág. 3 y 47.

Uranga, E. (20 de julio de 1968). Universidad Popular . La Prensa, pág. 3 y 35.

Uranga, E. (¿? de ¿? de 1968). Viaje del Canciller. La prensa, pág. 3 y 32.

Uranga, E. (1971). De Astucias Literarias. En Emilio, Astucias Literarias (págs. 31-35). México: Federación Editorial Mexicana.

Uranga, E. (1971). De Astucias Literarias. En E. Uranga, Astucias literarias (págs. 228-233). México: Federación Editorial Mexicana.

Uranga, E. (1971). Galeras de Astucias Literarias. En E. Uranga, Astucias Literarias (pág. 1971). México: Federación Editorial Mexicana .

Uranga, E. (1971). Prólogo de Astucias Literarias. México: Federación Editorial Mexicana .

Bibliografía

- Anónimo. (1968). *¡El Móndrigo! Bitácora del Consejo Nacional de Huelga*. México: Alba Roja.
- Bailey, R. W. (1979). *Authorship Attribution in a Forensic Setting*. Birmingham: AMLC.
- Belén, V. (1 de 08 de 2018). *En busca de... ¡El móndrigo! – el diario falso del movimiento estudiantil del 68*. Obtenido de Noticieros Televisa: <https://noticieros.televisa.com/especiales/mondrigo-diario-falso-movimiento-estudiantil-68/>
- Burrows, J. (2002). 'Delta': a Measure of Stylistic Difference and a Guide to Likely Authorship. *Literary and Linguistic Computing*, 267-287.
- Chung, C., & Chih-Jen, L. (2011). *LIBSVM -- A Library for Support Vector Machines*. Obtenido de Software available at <http://www.csie.ntu.edu.tw/~cjlin/libsvm>
- Daelemans, W. (2013). Explanation in Computational Stylometry. *Computational Linguistics and Intelligent Text Processing*, 451-462.
- Delgado, Á. (5 de octubre de 2003). Las confesiones de El Fish. *Proceso*, 1405, 22-25. Obtenido de <https://www.proceso.com.mx>.
- González, A. (05 de 01 de 2019). *¿Qué es Machine Learning?* Obtenido de CleverData: <https://cleverdata.io/que-es-machine-learning-big-data/>
- Grieve, J. (2007). Quantitative Authorship Attribution: An Evaluation of Techniques. *Literary and Linguistic Computing Advance Access*.
- Holmes, D. I. (1994). Authorship attribution. *Computers and the Humanities*, 87-106.
- Kilgariff, A. (2001). Comparing Corpora. *International Journal of Corpus Linguistics*, 97-133.

- Koppel , M., Schler, J., & Argamon, S. (2009). Computational Methods in Authorship Attribution. *Journal of the American Society for Information Science and Technology* , 9-26.
- López Escobedo, F., Solorzano Soto, J., & Sierra, G. (2016). Analysis of intertextual distances using multidimensional scaling in the context of authorship attribution. *Journal of Quantitative Linguistics*, 23(2), 154-176.
- López-Escobedo , F., Solorzano-Soto, J., & Sierra-Martínez, G. (2016). Analysis of Intertextual Distances Using Multidimensional Scaling in the Context of Authorship Attribution. *Jornal of Quantitative Linguistics* , 154-176.
- López-Escobedo, F., Sierra, G., & Solórzano, J. (2019). *SAUTEE: un recurso en línea para análisis estilométricos*. México: Grupo de Ingeniería Lingüística.
- Martré, G. (1998). *El movimiento popular estudiantil de 1968 en la novela mexicana*. México: UNAM.
- Monsiváis, C. (8 de 12 de 2018). *De libelos y libros*. Obtenido de <https://www.proceso.com.mx>: <https://www.proceso.com.mx/137972/de-libelos-y-libros>
- Munguía, J. R. (2016). *La otra guerra secreta: los archivos prohibidos de la prensa y el poder*. México : Random House.
- Padró, L. (2011). Analizadores Multilingües en FreeLing. *Linguamatica*, 3(2), 13-20.
- Sebastiani, F. (2002). Machine learning in automated text categorization. *ACM Computing Surveys* , 1-47.
- Sheridan, G. (04 de 12 de 2018). *Letras Libres* . Obtenido de Paz y Fuentes: cartas tlatelolcas (“el sector intelectual”): <https://www.letraslibres.com/mexico/literatura/paz-y-fuentes-cartas-tlatelolcas-el-sector-intelectual>

Stanczyk, U., & A. Cyran, K. (2008). Can punctuation marks be used as writer invariants? rough set-based approach to authorship attribution. *2nd EUROPEAN COMPUTING CONFERENCE*, 228-233.

Tasso, P. (04 de 21 de 2018). *La isla sin Este*. Obtenido de El mondriego, la novela de estado : <http://laislasineste-analisis.blogspot.com/2007/11/el-mndriego-la-novela-de-estado.html>

Valdés, G. M. (1972). *Operación 10 de Junio* . México : Ediciones Universo.