

Los lenguajes controlados y su aplicación a la traducción automática: 20 años después

Andrea Fernández Vivanco
José Manuel Bustos Gisbert



Andrea Fernández Vivanco
Universidad de Salamanca;
afernandezviv@usal.es;
ORCID: [0009-0000-0943-9763](https://orcid.org/0009-0000-0943-9763)



José Manuel Bustos
Gisbert
Universidad de Salamanca;
jbustos@usal.es;
ORCID: [0000-0001-9239-610X](https://orcid.org/0000-0001-9239-610X)

Resumen

Este estudio revisa el uso de lenguajes controlados como estrategia de preedición para mejorar la calidad de la traducción automática. Esbozamos un marco teórico del que se desprenden variables lingüísticas que sirven para diseñar un caso práctico. A partir de una evaluación cuantitativa y cualitativa, analizamos la productividad de la preedición y proponemos líneas de investigación futuras.

Palabras clave: lenguajes controlados, traducción automática, preedición, traducibilidad, calidad de la traducción automática.

Abstract

This study explores controlled languages as a pre-editing strategy to enhance machine translation quality. Drawing from a theoretical framework, we identify linguistic variables that shape the development of a practical case study. We evaluate the productivity of pre-editing through quantitative and qualitative analysis and propose avenues for future research.

Keywords: controlled languages, machine translation, pre-editing, translatability, machine translation quality.

Resum

Aquest estudi revisa l'ús de llenguatges controlats com a estratègia de preedició per millorar la qualitat de la traducció automàtica. Tracem un marc teòric del qual es desprenen variables lingüístiques que serveixen per dissenyar un cas pràctic. A partir d'una avaluació quantitativa i qualitativa, analitzem la productivitat de la preedició i proposem línies d'investigació futures.

Paraules clau: llenguatges controlats, traducció automàtica, preedició, traductibilitat, qualitat de la traducció automàtica.

Financiación

Este artículo se ha realizado en el marco del proyecto COMPLETEXT: Complejidad textual y lecturabilidad. Estudio aplicado a la IA generativa y la didáctica de las lenguas (Proyecto financiado por la JCyL FEDER-UE, SA067P24) y de un contrato predoctoral cofinanciado por la Junta de Castilla y León y el Fondo Social Europeo, al amparo de la Orden de 1 de diciembre de 2023, de la Consejería de Educación (BOCYL 233/2023).

1. Introducción: los lenguajes controlados más allá de la traducción

Kittredge (2014) define los lenguajes controlados (en adelante, LC) como *lenguajes contruidos* que, si bien no son producto de un proceso espontáneo, se basan en una lengua natural y restringen el uso en ciertos niveles lingüísticos, como el léxico, la sintaxis o la semántica. Puesto que los LC preservan la mayoría de las propiedades de esa lengua, sus hablantes pueden entender textos en LC de forma intuitiva (Kuhn, 2014). Recordemos al respecto que la facilidad de procesamiento y la consistencia del lenguaje fueron aspectos que llevaron a Chomsky (1993) a describir tres modelos de representación del lenguaje¹: en este contexto, los LC se pueden asimilar a los *lenguajes de estado finito* (en inglés, *Finite-State Languages*), en tanto que estos operan mediante un número determinado de reglas estrictas y simplifican las estructuras lingüísticas para la comunicación controlada².

La interdisciplinariedad de la bibliografía sobre los LC rinde cuenta del mosaico de agentes involucrados en su estudio, su desarrollo y su implementación. Las propuestas originadas en empresas como Ericsson, IBM o Airbus son ejemplo de la necesidad de armonizar flujos de comunicación empresarial (De Vecchi, 2018) y testigo de la extensión del inglés como *lingua franca* en compañías globalizadas. En ese contexto, las revisiones bibliográficas identifican varios tipos de LC auspiciados por organizaciones, como es el caso del ASD Simplified Technical English (2013), así como una cantidad considerable de LC promovidos por la investigación académica. En esta categoría se encuentran, entre otros, el OWL Simplified English (Cregan, Schwitter y Meyer, 2007), el E2V (Pratt-Hartmann, 2003) o el CLCM (Temnikova, 2012).

Si bien son varias las disciplinas que se ocupan del estudio de los LC, podemos destacar el interés que han suscitado para la producción científica en informática, lógica

¹ Recordemos que Chomsky distingue entre (a) finite-state languages, (b) phrase-structure grammars y (c) transformational grammars.

² Desde un punto de vista teórico, los LC se pueden diferenciar de propuestas afines, tales como los vocabularios controlados, las guías de estilo y los fragmentos del lenguaje. El vocabulario controlado alude a un subconjunto léxico limitado y codificado que puede formar parte de un LC, pero no constituye un LC por sí solo. Por su parte, una guía de estilo es un conjunto de recomendaciones sobre redacción, gramática y formato. Su aplicación puede ser más o menos sistemática y se diferencia de los LC en tanto que no persigue un propósito concreto. Finalmente, los fragmentos del lenguaje son una estrategia para descomponer empíricamente una lengua natural que tiene como objetivo analizar textos, pero no supone una intervención normativa o constructiva con un fin concreto.

formal, documentación y lingüística, donde incluimos los enfoques aplicados a la traducción. Los estudios en informática han propuesto LC como el CLOnE (Funk *et al.*, 2007) o el AIDA (Kuhn, Nagy y Krauthammer, 2013), mientras que la lógica formal ha dado lugar, entre otros, al CELT (Pease y Li, 2010) y a los silogismos de Sowa (2000). Los investigadores en documentación se han interesado principalmente por los LC desarrollados en empresas, en tanto que suponen un recurso para gestionar la información. Finalmente, los intentos de la lingüística por definir reglas de simplificación del lenguaje natural han dado lugar a LC como E-Prime (Bourland, 1965), pero también han incidido en disciplinas como la lectura fácil, el lenguaje claro y la simplificación textual automatizada (véase *infra*).

La pluralidad de LC, especialmente en lengua inglesa, motivó la clasificación PENS (Kuhn, 2014), que formula las nueve categorías contempladas en la Tabla 1. La clasificación evalúa el propósito de los LC (descrito por los códigos C, T, F), el canal (W, S), la autoría (A, I, G) y el dominio (D). Esta clasificación retoma los postulados de Schwitter (2002), que dividió los LC según su funcionalidad.

Code	Property
C	The goal is comprehensibility
T	The goal is translation
F	The goal is formal representation (including automatic execution)
W	The language is intended to be written
S	The language is intended to be spoken
D	The language is designed for a specific narrow domain
A	The language originated from academia
I	The language originated from industry
G	The language originated from a government

Tabla 1: Clasificación PENS de LC (Kuhn, 2014: 126)

Una de las principales variables (Huijsen, 1998) en su caracterización es la diferencia establecida entre los LC orientados a máquinas (MOCL, por sus siglas en inglés) y los orientados a humanos (HOCL). Los primeros persiguen programar sistemas de procesamiento del lenguaje natural que, una vez predefinidas ciertas reglas, identifiquen las infracciones del texto y sean capaces de aplicar las transformaciones necesarias. Los segundos privilegian la aplicación manual de las reglas, si bien el texto derivado puede ser procesado por una herramienta computacional.

Son varios los recursos tecnológicos que se basan en los MOLC, principalmente en el plano intralingüístico. La aplicación sistemática de restricciones para facilitar la comprensión de un texto ha dado lugar a correctores automáticos, a sugerencias de mejora de la redacción integrada en los procesadores de textos o a redactores asistidos como ArText (Da Cunha, 2022). Hasta ahora, estas aplicaciones se han servido principalmente del paradigma por reglas del procesamiento del lenguaje natural, si bien cabe esperar que los avances en arquitecturas *transformer* y modelos del lenguaje irrumpen pronto en su implementación.

En cualquier caso, podemos afirmar que las transformaciones textuales que se derivan de la aplicación de LC tienen que ver con tres dimensiones; a saber, la intralingüística, la interlingüística y la intersemiótica, como se muestra en la Ilustración 1.

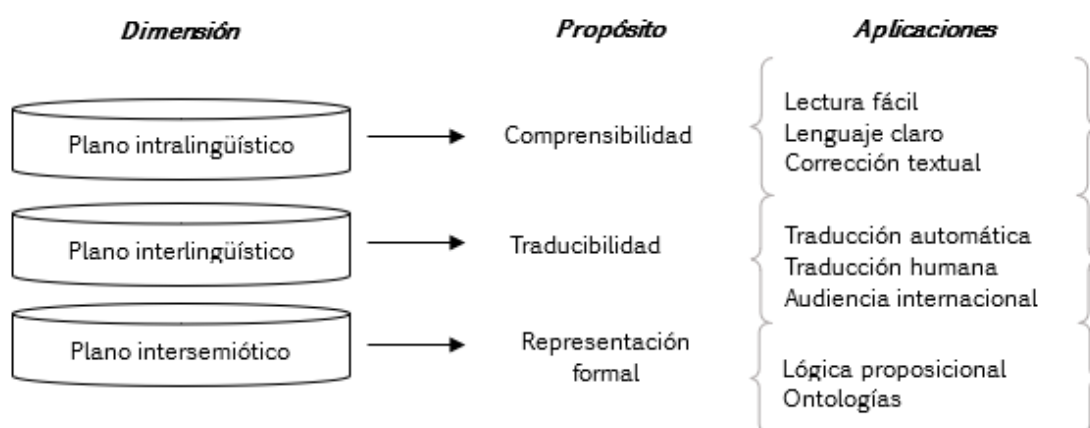


Ilustración 1: Dimensiones de los lenguajes controlados y aplicaciones relacionadas.

Nuestro estudio se centra en el plano interlingüístico, pues pretende evaluar las aplicaciones de los LC en la mejora de los resultados aportados por la TA. La clasificación PENS (Kuhn, 2014), presentada en la Tabla 1, organiza los LC entendidos como producto; es decir, contempla aquellas series definidas de reglas concebidas por un agente (académico, gubernamental o institucional) para su aplicación en un contexto comunicativo concreto. La gran mayoría de los LC que cumplen estas características, como los mencionados más arriba (ASD, OWL, CLOnE, etc.) han sido desarrollados en lengua inglesa y se aplican a una lengua de especialidad o un género textual concreto. Sánchez-Gijón y Kenny (2022) han señalado la preponderancia de los LC en contextos industriales y en la redacción de contenido destinado a una audiencia internacional mediante el empleo de una *lingua franca*. Esta realidad explica no solo que la mayoría de los LC se hayan desarrollado en ámbito anglohablante, sino también que la formulación de las directrices se inspire en las características lingüísticas y discursivas del inglés.

Los LC aplicados a la traducción han de observar la idiosincrasia del género textual y el sublenguaje empleados en lengua origen y lengua meta³. Sin embargo, la multimodalidad de los textos, la falta de equivalencia unívoca entre géneros textuales y la enormidad de los potenciales sublenguajes en el texto traducido nos hacen pensar que, en lugar de concebir los LC como *productos*, tal como Kuhn (2024) los presenta en su modelo PENS, la actividad traductora necesita abordarlos como *estrategia*. De ese modo, tal vez no se trate de definir las características propias de un único LC, sino más bien de fijar una serie de orientaciones que, según el contexto comunicativo, guíen el modo de actuar y faciliten la toma de decisiones. Es decir, en la preedición de textos que van a ser procesados con un motor de TA no se trataría tanto de aplicar *un* LC prefijado de manera general, sino de decidir qué conjunto de reglas propias de los LC se han de aplicar con el fin de aumentar la *traducibilidad* del texto. Esta selección consciente estaría orientada por el género textual, el motor empleado, la combinación lingüística, el ámbito de especialidad y determinados factores contextuales.

Considerando el plano intralingüístico que mencionamos en la Ilustración 1, observamos que ya se ha dotado a la lengua española con un corpus relativamente extenso de LC aplicados a la comprensibilidad (Seibel y Carlucci, 2021). Destacan los estudios en lectura fácil (García Muñoz, 2012; Jiménez Hurtado y Medina Reguera, 2022, entre otros), lenguaje claro (Real Academia Española, 2024; Montolío Durán, 2023; Da Cunha, 2022, entre otros) y simplificación textual (Saggion, 2015). Este grupo de LC ofrece ejemplos orientados a máquinas (Da Cunha, 2022, 2024) y otros pensados para la aplicación humana (Real Academia Española, 2024).

No obstante, la atención que han recibido los LC en el plano interlingüístico, orientados a mejorar la traducibilidad, parece más bien escasa y, sobre todo, se remonta a un estadio muy temprano de la TA. La propuesta de Ruiz Cascales y Sutcliffe (2003) fue uno de los primeros intentos de aplicar un conjunto de reglas para generar documentación técnica. Estos autores se propusieron trasponer al español las directrices formuladas por el AECMA Simplified English. Por su parte, Ramírez Polo (2012b) defendió una tesis doctoral que evaluaba el desempeño del LC en la TA centrándose en los entornos industriales. La autora propuso cuatro grandes bloques de normas orientados a reducir la ambigüedad léxica, sintáctica y contextual, así como controlar aspectos formales, como la puntuación y la ortografía.

Desde que se formularon estas propuestas, la TA ha progresado notablemente, evolucionando desde la basada en reglas (PBMT por sus siglas en inglés) hasta la integración de redes neuronales (Wu *et al.*, 2016) y modelos masivos del lenguaje (Schieferdecker, 2024). El objetivo principal de este estudio es comprobar la validez de las variables tradicionales de LC dado el estado de desarrollo de la TA, contemplando la aplicación de los LC como una estrategia de preedición. Para ello, partimos de la

³ Un sublenguaje es un subconjunto especializado de un lenguaje natural que se caracteriza por un vocabulario, una sintaxis y unos patrones de uso específicos adaptados a un dominio, una profesión o una comunidad discursiva específicos. Para una descripción completa véase Kittredge (1982:1-9).

relación de reglas de LC diseñada para la combinación lingüística español > inglés y elaborada a partir de Ramírez Polo (2012a) y O'Brien (2003), que esquematizamos en la Tabla 2:

X01	Evitar frases completas entre paréntesis
X02	Evitar errores ortográficos
X03	Evitar formas de plural entre paréntesis
X04	En una frase condicional, la condición ha de preceder a la acción
X05	Evitar oraciones pasivas sin agente explícito
X06	Evitar errores de puntuación
X07	Respetar el orden natural de los elementos sintácticos
X08	Acortar distancias entre sujeto y verbo
X09	Reducir el número de subordinadas por frase
X10	Evitar elipsis cuando el referente se encuentra en una frase diferente
X11	Evitar perífrasis verbales
X12	Evitar tiempos verbales compuestos
X13	Desarrollar siglas, acrónimos y abreviaturas
X14	Evitar el uso de modismos y expresiones locales
X15	Evitar que, en una oración, se repita una preposición para dos sintagmas preposicionales
X16	Evitar frases largas
X17	Evitar la doble negación
X18	Evitar el uso del gerundio

Tabla 2: Variables de LC aplicado a la TA español-inglés

Las adaptaciones lingüísticas promovidas por la aplicación de los LC tienen como objetivo simplificar el texto de partida, aumentar la consistencia, eliminar las construcciones gramaticales típicas de la lengua origen que generan problemas de traducción y desambiguar el contenido. No obstante, si analizamos detenidamente las reglas propuestas, podemos distinguir dos líneas diferentes asociadas a su diseño. Así, en algunos casos parece que el lenguaje se manipula con el objetivo de facilitar a la

herramienta de traducción automática (TA) la labor de descodificación del mensaje. Como resultado de ello, observamos que algunas reglas se orientan a eliminar construcciones polisémicas para reducir el número posible de significados identificables. Un ejemplo de esto es la regla X10: cuando el referente de una elipsis se encuentra en otra oración, a menudo la traducción hacia el inglés necesita de un pronombre personal que la máquina tiene que desambiguar (Guilou y Hardmeier, 2016). Sin embargo, en otros casos, observamos que ciertas reglas de LC persiguen facilitar la comprensión humana. Es lo que sucede con cuestiones como adelantar la prótasis en las oraciones condicionales (variable X04), que no tienen tanto que ver con el procesamiento computacional, sino más bien con facilitar la comprensión lectora del receptor. En ese sentido, son varias las investigaciones (Roturier, 2006; O'Brien, 2010) que demuestran que la aplicación de los LC mejora la lecturabilidad del texto.

2. Los lenguajes controlados en el flujo de trabajo de la traducción: la preedición

Los flujos de trabajo que incorporan herramientas de TA adoptan diversas formas e integran actividades tan variadas como la comunicación con el cliente o el desarrollo y la actualización de motores de traducción o de recursos léxicos y conceptuales. No obstante, se puede simplificar el proceso en tres grandes etapas: la preedición, el procesamiento del texto por un motor de TA y la posesición y/o revisión de la propuesta obtenida. La noción de preedición engloba todos los cambios realizados sobre el documento origen antes de procesarlo con el motor de TA, incluyendo tanto adaptaciones de formato como lingüísticas. Las estrategias adoptadas para adaptar el formato de un documento en la etapa de preedición incluyen, entre otras, la conversión a un archivo estandarizado de localización, la eliminación de imágenes o el reconocimiento óptico de caracteres. Por su parte, las adaptaciones lingüísticas consisten en modificar el texto origen para lograr un procesamiento óptimo del motor de TA y un resultado final que mejore la calidad y la idiomática, y reduzca los esfuerzos de posesición/revisión (O'Brien, 2006, 2010). Es precisamente en este escalón del flujo de traducción donde se aplican los LC, como se muestra en la Ilustración 2:

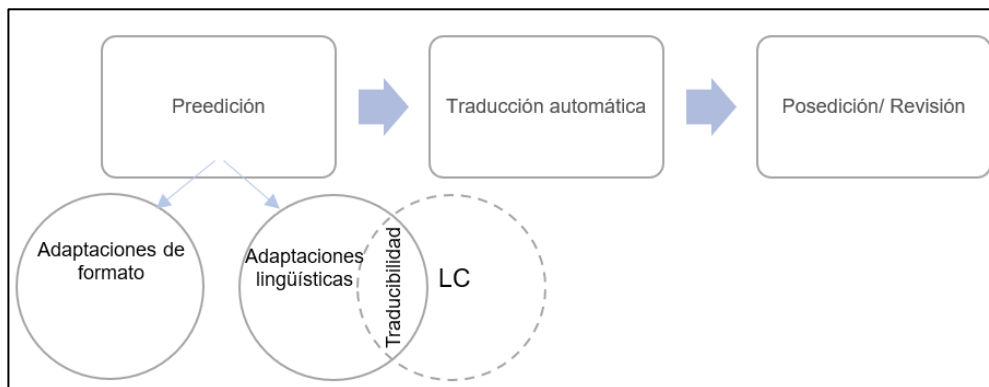


Ilustración 2: Situación de los lenguajes controlados en el flujo de trabajo de la TA.

3. Estudio de caso: descripción y metodología aplicada

Adaptar las estrategias de LC al estado de desarrollo tecnológico de los motores de TA es una empresa ambiciosa y requiere de un estudio exhaustivo que tome en consideración la diversidad de géneros, lenguas de especialidad, herramientas de TA y combinaciones lingüísticas que se pueden ver involucradas. Nuestro objetivo a tal efecto ha sido proponer un estudio de caso que sirva como preámbulo a investigaciones de mayor calado. En él, hemos aplicado las reglas de LC a un proceso de TA en la combinación español > inglés y hemos analizado su efecto en la calidad de la traducción proporcionada por el motor de traducción: para ello hemos procesado un mismo texto preeditado y sin preeditar con un motor de TA y hemos evaluado las potenciales mejoras de calidad.

Hemos seleccionado como texto origen un artículo periodístico publicado el 12 de noviembre de 2024 en *El País* (Pellicer, 2024), el diario generalista de mayor difusión en España. Dado que se trata de un estudio preliminar, hemos optado por este género al estar sociolectalmente marcado y no presentar de manera muy evidente rasgos prototípicos de la lengua de especialidad⁴. Es decir, nos permite describir un encargo de traducción ficticio pero plausible y, al mismo tiempo, soslayar errores de traducción derivados de la terminología propia de las lenguas de especialidad. Con el fin de abarcar la totalidad de las variables recogidas en la Tabla 2, incluidas aquellas que tienen que ver con errores sintácticos u ortotipográficos, hemos procesado tanto la noticia (823 palabras) como los doce primeros comentarios de usuarios de elpais.es (541 palabras). Recordemos que los comentarios de usuarios son un tipo textual que a menudo se traduce con sistemas de TA (cf. Lohar *et al.*, 2023).

En primer lugar, hemos segmentado el texto origen con la herramienta de traducción asistida por ordenador Trados Studio con el fin de etiquetar todas las infracciones de las reglas (X1-X16) en cada uno de los segmentos. En ese proceso hemos observado que algunas de las variables propuestas tienen una formulación interpretable, puesto que, tal y como indican Nyberg *et al.* (2003: 247), “las reglas de los LC orientados a máquinas suelen ser más específicas que las de LC orientados a humanos”⁵. Por ello, para las variables X9 y X16 se han generado etiquetas en tres subcategorías, lo que ha dado lugar a la clasificación que se muestra en la Tabla 3.

⁴ Cabe precisar que el texto periodístico se caracteriza por la naturaleza no especializada del destinatario, lo que condiciona la forma de escritura. Remitimos por su pertinencia a la noción de transparencia textual definida por Núñez Ladevéze (1993).

⁵ Traducción propia.

X09. Reducir el número de subordinadas por frase.		X16. Evitar frases largas.	
	X09A Nivel 2 del árbol de dependencia ⁶		X16A Más de 20 palabras
	X09B Nivel 3 del árbol de dependencia		X16B Más de 40 palabras
	X09C Nivel ≥ 4 árbol de dependencia		X16C Más de 60 palabras

Tabla 3: Etiquetas descriptoras asignadas a las variables X09 y X16

Los segmentos fueron editados manualmente a partir de las reglas descritas en la Tabla 2. Puesto que las modificaciones fueron incorporadas por los investigadores y no de manera automática, nuestro estudio de caso adopta una perspectiva HOCL. Fruto de la corrección llevada a cabo, el estudio ha permitido contemplar dos textos origen: uno sin preeditar (en adelante, TO) y otro preeditado (en adelante, TO-LC). El ejemplo (1) muestra un segmento en el que se observaron 5 infracciones a las reglas de LC y que fue editado para reducir la distancia entre sujeto y verbo (variable X08), disminuir el nivel de subordinación (variable X09C), eliminar perífrasis verbales (variable X11) y tiempos compuestos (variable X12), y reducir la longitud de la frase (variable X16B).

(1) TO: Lo sucedido en esa autonomía, donde se augura una ardua batalla política por la gestión de la dana –Compromís ha acusado al presidente, Carlos Mazón, de “homicidio imprudente”— llevó al resto de autonomías a apresurarse a actuar ante un fenómeno climático que, aunque se prevé más débil, ha demostrado que también puede ser impredecible.

TO- LC: En esa autonomía se augura una ardua batalla política por la gestión de la dana –Compromís ha acusado al presidente, Carlos Mazón, de “homicidio imprudente”. Lo sucedido allí llevó al resto de autonomías a actuar ante un fenómeno climático que ha demostrado también impredecibilidad, aunque se prevé más débil.

Nuestro estudio pretende comparar la calidad de dos textos meta: por un lado, la TA del texto origen sin preeditar (en adelante, TM) y por otro, la TA del preeditado (en adelante, TM-LC). Por ello, ambas versiones fueron procesadas a texto completo –es decir, sin segmentar– por el motor de TA DeepL Pro el 19 de noviembre de 2024 seleccionando la combinación Español > Inglés británico⁷. El corpus de análisis completo puede consultarse en Fernández Vivanco y Bustos Gisbert (2025). Ambas traducciones fueron evaluadas por seis informantes, todos ellos profesionales en ejercicio. El principal

⁶ La altura en el árbol de dependencia se ha calculado según Bohnet (2009): se toma el nodo incrustado en el nivel sintáctico más bajo de la oración y se mide la distancia entre ese nodo y la raíz del árbol (verbo principal del enunciado).

⁷ La selección del motor de traducción está motivada por la disponibilidad y el desempeño de DeepL Pro, que ha sido valorado como el motor de TA con mayor puntuación a partir de una evaluación LQA humana en la combinación inglés-español según Intento (2025: 55).

ámbito de especialidad de estos informantes es el institucional y los datos relacionados con su experiencia laboral, ocupación y empleador se desglosan en la Tabla 4:

Género	Combinación lingüística	País de residencia	Experiencia laboral	Ocupación	Empleador	Lengua materna	Institución
M	EN<>ES DE>EN	España	De 3 a 5 años	Intérprete y traductor	Sector público	EN, ES	CTIE ⁸ Ministerio del Interior
F	ES<>EN FR, RU>ES	España	De 5 a 10 años	Intérprete y traductora	Sector público	ES	CTIE Ministerio de Presidencia
F	ES<>EN FR, IT>ES	Irlanda	De 5 a 10 años	Traductora editorial	Cuenta propia	ES	
F	ES<>EN FR, PT>ES	Luxemburgo	De 5 a 10 años	Correctora de traducciones	Sector público	EN, ES	Tribunal de Justicia de la UE
F	ES, GL<>EN FR>ES, GL	Francia	De 5 a 10 años	Intérprete y traductora	Cuenta propia	ES, GL	
F	EN<>ES	España	De 5 a 10 años	Intérprete y traductora	Sector público	ES	CTIE Ministerio de Exteriores

Tabla 4: Datos sociodemográficos de los informantes

Como se muestra en la Ilustración 3, los seis evaluadores recibieron las dos traducciones. De ellos, los evaluadores pares tuvieron como punto de referencia el texto origen, mientras que la otra mitad recibió el texto origen preeditado. La decisión de dividir la muestra se fundamenta en el hecho de que los sistemas de TA suelen conservar la puntuación fuerte del TO. Si los evaluadores detectasen cambios en la longitud de los enunciados o en la distribución de los segmentos, podrían pensar que una de las traducciones es humana y la otra automática. Para reducir el posible sesgo derivado de ello, se decidió distribuir a la mitad de los evaluadores el TO y a la otra mitad el TO-LC.

⁸ Cuerpo de Traductores e Intérpretes del Estado.

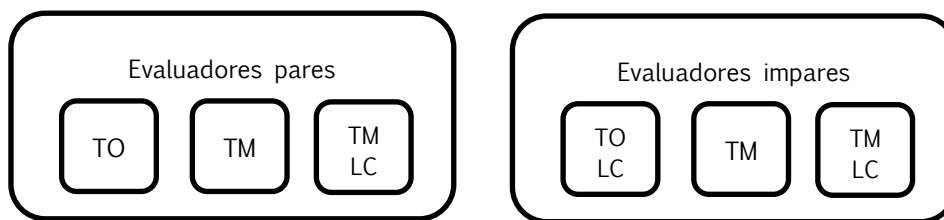


Ilustración 3: Esquema de la distribución de textos a los evaluadores

El fichero que recibieron los informantes estaba acompañado de las siguientes instrucciones:

En la hoja de cálculo adjunta encontrarás un texto periodístico sobre las inundaciones de la DANA de 2024 publicado en El País en octubre de 2024 (columna B) y dos propuestas de traducción (columnas C y G) para la edición en inglés del mismo periódico. Se incluyen también comentarios de lectores. En la primera columna encontrarás información contextual.

Para cada una de las propuestas de traducción distinguimos tres aspectos a evaluar: en azul, la precisión (contrasentidos, falsos sentidos, omisiones, adiciones, etc.); en verde, la fluidez (inconsistencias, registro, estilo), y en amarillo, la corrección (ortográfica y gramatical). Te pedimos que asignes una nota de 0 (puntuación más baja) a 5 (puntuación más alta) a cada unidad de traducción, es decir, a cada fila de la hoja de cálculo. Si el tiempo te lo permite, puedes describir los errores que hayan motivado tu puntuación con el mismo código de colores: contenido en azul, fluidez en verde y corrección en amarillo. Esta última parte es opcional, no queremos abusar de tu buena disposición y de tu tiempo.

Como se desprende de las instrucciones aportadas, se distinguieron tres aspectos de la traducción que los informantes tenían que calificar con una puntuación entre 0 y 5. Estos tres parámetros se han adaptado a partir de la categorización de errores del modelo MQM. La primera categoría, *precisión*, incluye errores como contrasentidos, falsos sentidos, no mismos sentidos, omisiones y adiciones. El segundo parámetro, *fluidez*, engloba las cuestiones relacionadas con la inconsistencia conceptual, el registro y el estilo. Finalmente, la tercera categoría se fija en la *corrección* ortográfica y gramatical. Si bien MQM ha actualizado recientemente estas categorías y ahora se refiere al parámetro *fluidez* como *convenciones lingüísticas*, nuestra adaptación pretende establecer un marco conciso y transparente para los evaluadores. La desviación media de los resultados aportados por los evaluadores que habían recibido el TO-LC con respecto a los que evaluaron el TO es del

0,12 %⁹. No obstante, el segmento 08 presentaba una desviación del 29 % debido a un cambio de sentido provocado por la preedición manual del texto origen. Este segmento fue invalidado y no se tuvo en cuenta en la evaluación de resultados. Asimismo, se ha decidido dejar de considerar el segmento 46, en el que se ha observado algún pequeño error de transcripción en el TO-LC. Por tanto, analizaremos en las secciones siguientes 52 segmentos.

4. Resultados, análisis y discusión

El análisis cuantitativo de las evaluaciones externas evidencia una mejora significativa de la calidad en todas las categorías de evaluación cuando el texto ha sido preeditado. En términos globales, la TA del texto sobre el que se han aplicado reglas de LC ha recibido una valoración general que es un 18 % superior.

4.1. Errores de contenido

Se ha observado una mejora más notable en la categoría de evaluación relacionada con la precisión conceptual, donde el TM-LC ha alcanzado una puntuación un 22 % superior al TM. En ningún segmento la puntuación de la categoría de precisión disminuye, es decir, no se han inducido errores derivados de la aplicación de LC.

Es relevante destacar que, al procesar el texto origen sin preeditar, el motor de TA arroja errores de contenido graves o muy graves en forma de contrasentidos, falsos sentidos u omisiones, por lo que varios segmentos han recibido puntuaciones muy bajas en ese apartado (entre 0 y 1). Especialmente paradigmático es el segmento recogido en el ejemplo (2):

- (2) TO: España se enfrenta a otra dana solo dos semanas después de la catástrofe de Valencia
TO-LC: España se enfrenta a otra depresión aislada en niveles altos (DANA) solo dos semanas después de la catástrofe de Valencia
TM: Spain faces another drought just two weeks after Valencia disaster
TM-LC: Spain faces another isolated depression at high levels (DANA) only two weeks after the catastrophe in Valencia.

En este ejemplo la puntuación promedio del TM para la categoría de contenido es de 0, mientras que en TM-LC asciende a 4,2. Esto se debe a que el motor de TA arroja *drought* (sequía) como equivalente de *dana*, introduciendo así un contrasentido.

⁹ La desviación hace referencia a la diferencia porcentual entre la puntuación media de los evaluadores que recibieron el TO y quienes recibieron el TO-PT.

Otra de las variables relacionadas con la mejora de la precisión conceptual corresponde a la eliminación de elipsis cuando el referente se encuentra en otra oración, tal y como se muestra en el ejemplo (3)¹⁰.

(3) TO: Pasadas las 21.00, se dio la alerta roja para Málaga, y sobre las 23.00, para el litoral sur de Tarragona. Para el resto, los avisos eran naranja y amarillos.

TO-LC: Las autoridades dieron la alerta roja pasadas las 21.00 para Málaga, y sobre las 23.00, para el litoral sur de Tarragona. Los avisos eran naranjas y amarillos para el resto de las regiones.

TM: After 21:00, a red alert was issued for Malaga, and at around 23:00 for the southern coast of Tarragona. For the rest, the warnings were orange and yellow.

TM-LC: The authorities issued a red alert after 21:00 for Malaga, and at around 23:00 for the southern coast of Tarragona. The warnings were orange and yellow for the rest of the regions.

4.2. Errores de corrección

En el parámetro corrección, la puntuación de la traducción tras aplicar LC aumenta en un 18 %. Cabe mencionar que las dos variables relacionadas tradicionalmente con la corrección –evitar errores ortográficos (X02) y de puntuación (X06)– no introducen cambios significativos de calidad. Sin embargo, destaca la productividad de la variable X07, que obliga a respetar el orden natural de los elementos de la oración, así como de la variable X10, que proscribía las elipsis cuando el referente se encuentra en una oración diferente. Obsérvense los ejemplos (4) y (5), que obtienen una mejora en el apartado de corrección del 83 % y el 117 % respectivamente.

(4) TO: Y todo ese nuevo capítulo de temor se producía justo en el arranque de la cumbre del clima, que en esta ocasión se celebra en Bakú (Azerbaián). Hasta allí llevó el presidente del Gobierno, Pedro Sánchez, la tragedia de Valencia.

TO-LC: Y todo ese nuevo capítulo de temor se producía justo en el arranque de la cumbre del clima. En esta ocasión la cumbre se celebra en Bakú (Azerbaián). El presidente del Gobierno, Pedro Sánchez, llevó hasta allí la tragedia de Valencia.

TM: And all this new chapter of fear came just at the start of the climate summit, which on this occasion is being held in Baku (Azerbaijan). It was there that the president of the government, Pedro Sánchez, took the tragedy of Valencia.

TM-LC: And all this new chapter of fear came just at the start of the climate summit. This time the summit is being held in Baku (Azerbaijan). The president of the government, Pedro Sánchez, brought the Valencia tragedy there.

¹⁰ El subrayado es nuestro.

(5) TO: En Valencia, que no se preocupen, tienen a Mazón y al pp¹¹

TO-LC: L En Valencia, que no se preocupen, tienen a Mazón y al Partido Popular.

TM: In Valencia, don't worry, they have Mazón and the PP.

TM-LC: In Valencia, don't worry, they have Mazón and the Partido Popular.

Sobre el ejemplo (5) hemos de apostillar que no solo se ha registrado una mejora en la categoría de corrección, sino en las tres categorías. Parece, por tanto, que en este caso el cambio derivado de desarrollar la sigla genera una traducción más idiomática y más precisa en cuanto al contenido.

4.3. Errores de fluidez

La categoría sometida a evaluación que presenta una mejora de la calidad más moderada es la de *fluidez*, pues se sitúa en el 12 %. Destaca la mejora asociada a evitar construcciones de verbo más infinitivo con arreglo a la variable X11, como se muestra en el ejemplo (6).

(6) TO: El Gobierno catalán de Salvador Illa, además, decidió sacar de la agenda de los hospitales las visitas no urgentes en los centros de las zonas amenazadas.

TO-LC: El Gobierno catalán de Salvador Illa, además, sacó de la agenda de los hospitales las visitas no urgentes en los centros de las zonas amenazadas.

TM: The Catalan government of Salvador Illa also decided to remove non-urgent visits to hospitals in the threatened areas from the agenda.

TM-LC: The Catalan government of Salvador Illa also removed non-emergency visits to hospitals in the threatened areas from the schedule.

Sobre la variable X11 interesa destacar que se han registrado mayores porcentajes de mejora cuando la construcción de verbo más infinitivo utiliza verbos no modales, como en el ejemplo recogido. Por su parte, las construcciones perifrásticas con valor modal no han generado errores en la TA.

La categoría de fluidez es la única que registra (muy excepcionalmente, eso sí) alguna disminución de la calidad para ciertos segmentos, como observamos en (2) con el caso del acrónimo *dana*¹². Importa ahora recordar que en la edición del artículo se aplicaron las reglas de LC de forma metódica y sin atender a factores contextuales. En otras palabras, si bien los cambios de preedición se realizaron manualmente, se efectuó una aplicación sistemática. Dado que la variable X13 indica que se han de desarrollar las

¹¹ El ejemplo (5) no ha sido extraído de la noticia, sino de un comentario en elpais.es.

¹² El lema *dana* fue incluido en el Diccionario de la Lengua Española en diciembre de 2024. Puesto que la presente investigación se desarrolló durante noviembre de 2024, se trata como acrónimo. A fecha de 27 de febrero de 2025, la traducción que ofrece DeepL Pro para el TO del ejemplo (2) sigue siendo: "Spain faces another drought just two weeks after Valencia disaster".

siglas, el TO-LC del ejemplo (2) incorpora un compuesto nominal terminológico en el título de una noticia, algo discordante con las convenciones del género y el estilo periodístico. Como consecuencia de ello, la puntuación obtenida en el parámetro fluidez desciende. Efectivamente, cuando una sigla se repite varias veces a lo largo del texto y se desarrolla en todos los casos, los evaluadores han percibido una disminución notable de fluidez. También se valora negativamente desarrollar aquellos acrónimos utilizados para designar una institución, como se muestra en (7).

(7) TO: La Aemet activa la alerta roja en Málaga y en el litoral sur de Tarragona.

TO-LC: La Agencia Española de Meteorología (Aemet) activa la alerta roja en Málaga y en el litoral sur de Tarragona.

TM: The Aemet activates a red alert in Malaga and the southern coast of Tarragona.

TM-LC: The Spanish Meteorological Agency (AEMET) activates a red alert in Malaga and the southern coast of Tarragona.

No obstante, incluso en el caso de los acrónimos que no actúan como nombres propios se observa una puntuación más alta en la categoría precisión. En (8) se muestra un ejemplo más de cómo una sigla genera problemas graves de contenido en la TA.

(8) TO: Para el que no termino la egb Y por qué tiene que dimitir el presidente del gobierno?¹³

TO-LC: Para el que no termino la educación general básica: ¿Y por qué tiene que dimitir el presidente del gobierno?

TM: And why does the president of the government have to resign? (sic)

TM-LC: For those of you who have not finished your basic general education: Why does the president of the government have to resign?

Del ejemplo (8) llama la atención que, cuando el texto está sin preeditar, el motor de TA renuncia a traducir un fragmento de siete palabras del texto original. En una comprobación posterior advertimos que, si procesamos el texto completo con el motor DeepL Pro, se omite una posible traducción para el sintagma “para el que no termino la egb”. No obstante, si se procesa únicamente el segmento en cuestión, el mismo motor proporciona la siguiente traducción: “For those who didn't finish the egb And why does the president of the government have to resign?”. Esta paradoja tal vez se pueda explicar por el hecho de que DeepL ha incorporado de manera reciente un módulo denominado LLM NextGen (DeepL, 2024); pues bien, algunas evaluaciones iniciales sobre el desempeño de los modelos masivos del lenguaje aplicados a la TA han demostrado que los errores relacionados con omisiones o con información no traducida tienen una mayor incidencia en la TA con LLM (Feng *et al.*, 2024).

¹³ El ejemplo (8) no ha sido extraído de la noticia, sino de un comentario en elpais.es.

Entre los segmentos que obtienen una peor valoración para la traducción del texto preeditado en el parámetro de fluidez se encuentran también varias oraciones que habían sido segmentadas con arreglo a la variable X16 (evitar frases largas). Así, los evaluadores valoraron negativamente fragmentar las oraciones con una longitud entre 20 y 30 palabras, como se muestra en el ejemplo (9), que segmenta una oración de 25 palabras en dos oraciones de 14 y 15 palabras respectivamente.

(9) TO: Las autoridades decidieron cerrar multitud de escuelas, pidieron no salir a la carretera e intentaron protegerse otra vez de posibles inundaciones de arroyos y barrancos.

TO-LC: Las autoridades cerraron multitud de escuelas y ellas pidieron no salir a la carretera. Ellas intentaron otra vez la protección ante posibles inundaciones de arroyos y barrancos.

TM: The authorities decided to close many schools, asked not to go out on the road and tried again to protect against possible flooding of streams and ravines.

TM-LC: The authorities closed a lot of schools and they asked not to go out on the road. They tried again to protect against possible flooding of streams and ravines.

4.4. Sobre las variables aplicadas

En primer lugar, los datos de rendimiento de las variables de preedición confirman la pertinencia de desglosar la regla 16 (evitar frases largas) en tres según la extensión de los segmentos. De hecho, esa división permite observar que la aplicación de la X16A (frases de más de 20 palabras) casi triplica la de X16C (frases de más de 60 palabras). Sin embargo, el desglose de la variable 9, ligada al nivel de subordinación, no parece necesario, pues las tres opciones resultantes tienen un rendimiento desagregado bastante limitado. Por ello, para el análisis de resultados que se ofrece a continuación, se sumarán los resultados correspondientes a X9A, X9B y X9C en una única regla X9. El rendimiento de las veinte reglas aplicadas en el proceso de edición es muy desigual, como se observa en la Tabla 5:

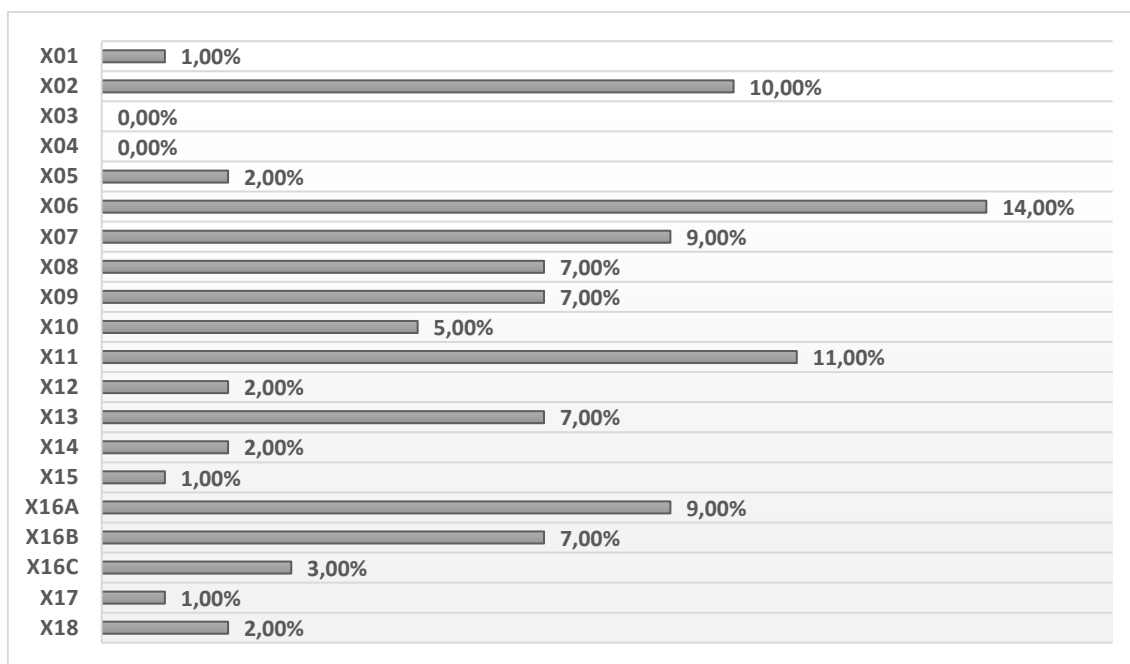


Tabla 5. Rendimiento de las normas de preedición

Cabe destacar, en primer lugar, que dos de ellas no se aplicaron en ninguna ocasión:

X03. Evitar formas de plural entre paréntesis

X04. En una frase condicional, la condición ha de preceder a la acción

Y, además, siete reglas tienen un rendimiento inferior al 2 %, por lo que se aplican solamente en una o dos ocasiones:

X01 Evitar frases completas entre paréntesis	1,00 %
X15 Evitar que, en una oración, se repita una preposición para dos sintagmas preposicionales	1,00 %
X17 Evitar la doble negación	1,00 %
X05 Evitar oraciones pasivas sin agente explícito	2,00 %
X12 Evitar tiempos verbales compuestos	2,00 %
X14 Evitar el uso de modismos y expresiones locales	2,00 %
X18 Evitar el uso del gerundio	2,00 %

Son, pues, nueve las reglas cuya pertinencia en el proceso de preedición de los textos originales se debe poner al menos en duda y someter a una valoración asentada en un corpus mayor que aumente la fiabilidad de los resultados.

En el extremo opuesto, se debe destacar en términos de cobertura que nueve reglas (todas ellas con un rendimiento superior al 7 %) dan cuenta de algo más del 80,81 % de las operaciones efectuadas¹⁴:

X06. Evitar errores de puntuación	14 %
X11. Evitar perífrasis verbales	11 %
X02. Evitar errores ortográficos	9 %
X07. Respetar el orden natural de los elementos sintácticos	9 %
X16A. Evitar frases largas de más de 20 palabras	9 %
X08. Acortar distancias entre sujeto y verbo	7 %
X09. Reducir el número de subordinadas por frase	7 %
X13. Desarrollar siglas, acrónimos y abreviaturas	7 %
X16B. Evitar frases largas de más de 40 palabras	7 %

El número de reglas aplicadas a cada segmento del corpus de trabajo oscila entre 0 y 6; el promedio es de 1,92 y la moda es de 1. La distribución se recoge en la Tabla 6:

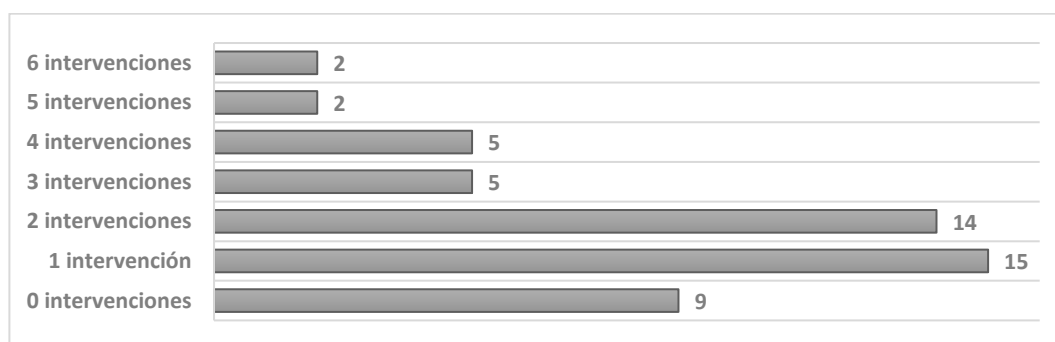


Tabla 6. Rendimiento de las normas de preedición

Existe cierta correlación entre el número de intervenciones y la extensión de los segmentos, tal y como se muestra en la Tabla 7:

¹⁴ En este sentido, los resultados se pueden asimilar a los hallazgos de Ramírez Polo (2012a).

<i>Extensión</i>	<i>N.º segmentos</i>	<i>% Segmentos</i>	<i>N.º Intervenciones</i>	<i>% Intervenciones</i>	<i>Promedio</i>
51+	6	11,54 %	28	28 %	4,67
40 a 50	3	5,77 %	10	10 %	3,33
30 a 40	4	7,69 %	9	9 %	2,25
20 a 30	14	26,92 %	22	22 %	1,47
10 a 20	18	34,62 %	21	21 %	1,11
0 a 10	7	13,46 %	10	10 %	1,43

Tabla 7. Intervenciones y extensión de los segmentos

El bloque correspondiente a los segmentos con una extensión entre 10 y 20 palabras es el mayor (35,19 %), pero solo suma un 20,20 % de las intervenciones. Como resultado de ello el promedio desciende a 1,05. En cambio, los segmentos con una extensión superior a 50 palabras suponen solo un 11,11 % del corpus; sin embargo, acumulan más de un 28 % de las intervenciones, por lo que el promedio se dispara en relación con los resultados globales del análisis. En términos cualitativos cabe destacar dos cosas: por un lado, que la variable X16 destinada a reducir la extensión de los segmentos se aplica, como es lógico, en los seis segmentos; y, por otro, que la variable X09, orientada a reducir el número de subordinadas, se aplica en cinco de los seis segmentos. Parece, por tanto, que serían dos reglas de naturaleza primaria en el proceso de preedición¹⁵.

6.5. Sobre la mejora de la traducción aportada

Como ya se apuntó antes, la preedición supone un promedio general de mejora del 18,17 % en el TM-LC respecto de la correspondiente al TM: en 32 segmentos (61,54 %) se observa algún porcentaje de mejora la traducción; en ellos, el promedio es del 31,23 %. En 13 (25 %) la traducción no mejora ni empeora. Y en un 7 (13,46 %) la traducción empeora tras la preedición del texto original; en estos, el promedio es del -7,78 %.

El porcentaje de mejora se mueve en un rango del 94,44 % al -16,67 %. Los segmentos se pueden clasificar en 6 categorías, que resumimos en la Tabla 8:

¹⁵ En relación con todo lo dicho, parece oportuno recordar ahora que Bustos Gisbert (2017: 92 y 93) ha demostrado, a partir de un corpus de 369 textos (15 000 palabras) divididos en 9 400 segmentos/enunciados, que en el texto escrito en español de naturaleza general “el promedio de extensión de los enunciados se fija en torno a 17 palabras”, y que “el 90 % de las unidades se mueven en una franja entre 6 y 30”. Asimismo, ha comprobado que “el 67 % de los segmentos se construye sintácticamente como oraciones compuestas, mientras que el 31,5 % responde a la estructura de una oración simple”; de los primeros, “el 50 % de los casos incluye una sola predicación subordinada, en el 30 %, incorpora 2 y en el 13,5 % son tres. Las unidades con cuatro o más predicaciones suman algo menos del 7 %”.

<i>Porcentaje de mejora</i>	<i>N.º de segmentos</i>	<i>% segmentos</i>
<i>Más del 61 %</i>	6	11,54 %
<i>De 41 a 60 %</i>	6	11,54 %
<i>De 21 a 40 %</i>	5	9,62 %
<i>De 1 a 20 %</i>	15	28,85 %
<i>0 %</i>	13	25,00 %
<i>De -1 a -17 %</i>	7	13,46 %

Tabla 8. Porcentajes de mejora

Valoraremos en primer lugar los 17 segmentos (32 %) en los que el porcentaje es superior al 20 %. Los datos se recogen en la Tabla 9:

<i>Segmento</i>	<i>LG⁶</i>	<i>Intervenciones</i>	<i>Reglas aplicadas</i>	<i>Valoración TM</i>	<i>Valoración TM-LC</i>	<i>Mejora (%)</i>
<i>Seg_34</i>	9	2	X02, X06	2,78	4,67	94,44 %
<i>Seg_01</i>	15	1	X13	2,83	4,22	69,44 %
<i>Seg_31</i>	20	1	X02	3,67	4,94	63,89 %
<i>Seg_35</i>	55	5	X02, X06, X07, X13, X16B	3,11	4,39	63,89 %
<i>Seg_23</i>	26	2	X11, X16A	2,50	3,78	63,89 %
<i>Seg_06</i>	62	4	X09, X12, X13, X16C	2,33	3,56	61,11 %
<i>Seg_41</i>	17	3	X02, X06, X13	3,72	4,83	55,56 %
<i>Seg_38</i>	2	2	X02, X06	2,13	3,13	50,00 %
<i>Seg_16</i>	25	1	X10	3,78	4,67	44,44 %
<i>Seg_26</i>	27	3	X06, X09, X11	3,83	4,72	44,44 %
<i>Seg_30</i>	4	1	X10	3,72	4,61	44,44 %
<i>Seg_10</i>	42	2	X08, X16B	3,11	3,94	41,67 %
<i>Seg_14</i>	9	2	X07, X10	4,11	4,83	36,11 %
<i>Seg_04</i>	15	1	X13	3,94	4,61	33,33 %
<i>Seg_44</i>	44	2	X06, X16A	3,44	4,06	30,56 %
<i>Seg_15</i>	55	4	X07, X09, X15, X16B	3,78	4,27	24,44 %
<i>Seg_28</i>	22	1	X18	4,50	4,94	22,22 %

Tabla 9. Segmentos con una mejora superior al 20%

¹⁶ Recuento de palabras.

Tal mejora descansa en la aplicación de 13 reglas:

X06. Evitar errores de puntuación	6
X13. Desarrollar siglas, acrónimos y abreviaturas	5
X02 Evitar errores ortográficos	5
X10 Evitar elipsis cuando el referente se encuentra en una frase diferente	3
X16A. Evitar frases largas de más de 20 palabras	3
X16B. Evitar frases largas de más de 40 palabras	3
X07. Respetar el orden natural de los elementos sintácticos	3
X09. Reducir el número de subordinadas por frase	3
X11. Evitar perífrasis verbales	2
X12. Evitar tiempos verbales compuestos	1
X15. Evitar que, en una oración, se repita una preposición para dos sintagmas preposicionales	1
X18. Evitar el uso del gerundio	1
X08. Acortar distancias entre sujeto y verbo	1

Un mayor número de reglas aplicadas no genera, necesariamente, un porcentaje de mejora mayor. De hecho, el promedio de intervenciones en este primer grupo es de 2,17 y la moda es de 2. Debemos destacar dos cuestiones fundamentales. La primera concierne a las variables más productivas, es decir, aquellas que se repiten en los segmentos que han registrado un mayor porcentaje de mejora; son las siguientes:

X06. Evitar errores de puntuación. (6 veces)
X16A+X16B. Evitar frases largas de más de 20/40 palabras (3+3 veces)
X13. Desarrollar siglas, acrónimos y abreviaturas (5 veces)
X02 Evitar errores ortográficos (5 veces)

La segunda tiene que ver con aquellas normas que, aun cuando no se repiten en demasiados segmentos, sí que conllevan un porcentaje de mejora elevado, pues se aplican sin combinar con otras normas:

X10 Evitar elipsis cuando el referente se encuentra en una frase diferente

X18 Evitar el uso del gerundio

La segunda categoría de análisis que hemos de considerar es el de aquellos segmentos en los que la preedición empeora la calidad de la traducción (13,46 %). Los recogemos en la Tabla 10:

<i>Segmento</i>	<i>LG</i>	<i>Intervenciones</i>	<i>Reglas aplicadas</i>	<i>Valoración TM</i>	<i>Valoración TM-LC</i>	<i>Mejora (%)</i>
<i>Seg_07</i>	27	4	X07, X08, X12, X16C	4,56	4,22	-16,67 %
<i>Seg_09</i>	25	2	X11, X16A	4,50	4,28	-11,11 %
<i>Seg_36</i>	19	1	X06	3,80	3,60	-10,00 %
<i>Seg_11</i>	68	4	X07, X08, X09	4,17	4,06	-5,56 %
<i>Seg_24</i>	27	1	X16A	4,67	4,56	-5,56 %
<i>Seg_20</i>	34	3	X08, X11, X16A	4,56	4,50	-2,78 %
<i>Seg_27</i>	13	2	X10, X11	4,72	4,67	-2,78 %

Tabla 10. Segmentos con empeoramiento de la calidad

Solo en los tres primeros segmentos se observa un descenso de cierta relevancia y no parece atribuible de manera evidente a la aplicación de alguna de las normas de preedición empleadas. De hecho, como hemos visto más arriba la aplicación de las reglas X06, X08, X11, X12 y X16 pueden incidir muy positivamente en la calidad de la TA. En todos los casos, se debe destacar que el descenso de la calidad se observa normalmente asociada a la pérdida de fluidez de la traducción. En ellos, como se observa en la Tabla 11, la calificación referida a ella desciende:

<i>Segmento (Fluidez: F)</i>	<i>Valoración TM</i>	<i>Valoración TM-LC</i>	<i>Mejora de TM-LC %</i>
<i>F_seg_07</i>	4,33	3,33	-50,00 %
<i>F_seg_09</i>	4,50	4,00	-25,00 %
<i>F_seg_11</i>	4,17	3,67	-25,00 %
<i>F_seg_20</i>	4,33	4,17	-8,33 %
<i>F_seg_24</i>	4,00	3,67	-16,67 %
<i>F_seg_27</i>	5,00	4,67	-16,67 %
<i>F_seg_36</i>	3,00	2,80	-10,00 %

Tabla 11. Segmentos con empeoramiento de calidad: la valoración de la fluidez

Es importante subrayar que el número de valoraciones negativas en los textos preeditados correspondientes al apartado de fluidez (11 segmentos) dobla a la suma de las correspondientes a corrección (5 segmentos) y a precisión (6 segmentos). Lo resumimos en la Tabla 12:

<i>Segmento (Fluidez: F; Corrección: C; Precisión: P)</i>	<i>Valoración TM</i>	<i>Valoración TM-LC</i>	<i>Mejora de TM-LC %</i>
<i>C_seg_09</i>	4,33	4,17	-8,33 %
<i>C_seg_10</i>	4,67	4,33	-16,67 %
<i>C_seg_20</i>	4,67	4,50	-8,33 %
<i>C_seg_32</i>	4,00	3,83	-8,33 %
<i>C_seg_36</i>	4,80	4,00	-40,00 %
<i>F_seg_01</i>	4,33	3,17	-58,33 %
<i>F_seg_07</i>	4,33	3,33	-50,00 %
<i>F_seg_09</i>	4,50	4,00	-25,00 %
<i>F_seg_10</i>	4,17	3,83	-16,67 %
<i>F_seg_11</i>	4,17	3,67	-25,00 %
<i>F_seg_12</i>	4,50	3,67	-41,67 %
<i>F_seg_20</i>	4,33	4,17	-8,33 %
<i>F_seg_24</i>	4,00	3,67	-16,67 %
<i>F_seg_27</i>	5,00	4,67	-16,67 %
<i>F_seg_32</i>	3,50	3,33	-8,33 %
<i>F_seg_36</i>	3,00	2,80	-10,00 %
<i>P_seg_07</i>	4,50	4,33	-8,33 %
<i>P_seg_08</i>	4,67	4,00	-33,33 %
<i>P_seg_15</i>	4,50	4,00	-25,00 %
<i>P_seg_17</i>	4,50	4,00	-25,00 %
<i>P_seg_19</i>	4,67	4,50	-8,33 %
<i>P_seg_46</i>	4,83	4,67	-8,33 %

Tabla 12. Segmentos con descenso de la valoración en corrección, fluidez y precisión

Por último, fijaremos la atención en los 13 segmentos (25 %) en los que la preedición no supone cambio en la valoración. Recogemos en la Tabla 13 los cuatro en los que se ha aplicado alguna norma de preedición:

<i>Segmento</i>	<i>LG</i>	<i>Intervenciones</i>	<i>Reglas aplicadas</i>	<i>Valoración TM</i>	<i>Valoración TM-LC</i>	<i>Mejora (%)</i>
<i>Seg_12</i>	18	2	X05, X11	4,39	4,39	0,00 %
<i>Seg_22</i>	26	1	X11	3,89	3,89	0,00 %
<i>Seg_37</i>	8	1	X06	5,00	5,00	0,00 %
<i>Seg_39</i>	5	2	X02, X06	5,00	5,00	0,00 %

Tabla 13. Segmentos con mejora 0

Cabe destacar dos cuestiones; de un lado, el valor 0 se asocia a segmentos de extensión reducida: tres de los valorados son inferiores a 20 palabras; de otro lado, llama la atención que la regla 11 (Evitar perífrasis verbales) tiene alguna relevancia en este contexto de mejora 0 y, sin embargo, es la segunda que más se aplica en todo el texto.

Con todo, observamos que algunos segmentos incumplen varias reglas y, por tanto, sufren distintas modificaciones. En estos casos, no es posible saber cuál es la responsable de la mejora o el empeoramiento.

7. Conclusiones

El análisis expuesto en las páginas anteriores apunta a que, de manera general, la preedición tiene un efecto positivo en la respuesta que nos aportan los motores de TA. Solo en el caso concreto de la fluidez de los textos se puede observar algún efecto negativo que, no obstante, es siempre porcentualmente muy inferior si se compara con la mejora global observada.

No obstante, es necesario avanzar en lo que tiene que ver con la aplicación de los LC a motores de traducción automática neuronal, pues las normas con las que ahora contamos no son en todos los casos válidas ni fiables: las más aplicables no generan necesariamente mejoras mayores; y, a la inversa, algunas de las menos usadas pueden tener efectos cualitativamente mayores en el resultado propuesto por el motor. El motivo es obvio: las propuestas existentes estaban pensadas para motores de naturaleza estadística (PBMT), que afrontaban dificultades distintas a las que se enfrenta un motor construido sobre redes neuronales.

Se hace, por tanto, necesario reformular la taxonomía de reglas aplicables, pero entendiendo siempre que no se trata tanto de cambiar la forma en la que se redacten los textos originales como de proponer protocolos de intervención en los textos antes de que sean tratados por un motor de TA o por una herramienta de tratamiento de

lenguaje del tipo de los GPT. En esa dirección, una opción puede ser mejorar las reglas aplicando los avances observados en los estudios de lecturabilidad y proponer reglas asociadas no tanto al concepto de *lenguaje controlado* como al de *lenguaje claro*¹⁷. Desde esa perspectiva, podremos fijar respuestas a cuestiones referidas a los aspectos propios de la construcción de los enunciados que afectan al nivel de claridad de los mismos.

Quedan, no obstante, algunas cuentas pendientes a las que no hemos podido ni siquiera acercarnos en estas páginas. En primer lugar, la necesidad de replicar análisis como el presentado en el contexto de diferentes ámbitos de especialización, lo que supone, necesariamente, fijar un modelo consistente de análisis en el nivel léxico, muy especialmente en lo referido a las terminologías empleadas. En segundo lugar, fijar los límites en la aplicación de las reglas en lo que tiene que ver con las convenciones de los géneros textuales: de hecho, se está observando que, en comunidades discursivas, como es el caso de las del español y el inglés, la aplicación de principios del lenguaje claro a los textos está provocando cambios, y no precisamente menores, en las convenciones de género (Williams, 2023; Pistola y Viñuales-Ferreiro, 2021). Por último, habrá que estudiar vías que exploren las posibles automatizaciones del proceso de preedición. Es nuestro propósito acometer en el futuro algunas de estas cuestiones, orientando nuestras investigaciones hacia modelos que integren la innovación tecnológica con el control humano.

9. Bibliografía

- ASD Aerospace and Defence Industries Association of Europe (2013). *Simplified Technical English. Specification ASD-STE100*, v. 6.
- Bohnet, Bernd (2009). Efficient Parsing of Syntactic and Semantic Dependency Structures. In: *Proceedings of the Thirteenth Conference on Computational Natural Language Learning (CoNLL 2009): Shared Task*. Association for Computational Linguistics, pp. 67–72. <<https://aclanthology.org/W09-1210/>>. [Accessed: 20250428].
- Bourland, David (1965). A linguistic note: Writing in E-prime. *General Semantics Bulletin*, v. 32, n. 3, pp. 111–114.
- Bustos Gisbert, José M. (2017). Naturaleza sintáctica de los enunciados textuales en el discurso escrito. *ELUA*, v. 31, pp. 67–95. <<https://doi.org/10.14198/ELUA2017.31.04>>. [Accessed: 20250428].
- Chomsky, Noam (1956). Three models for the description of language. *IRE Transactions on Information Theory*, v. 2, n. 3, pp. 113–124. <<https://doi.org/10.1109/TIT.1956.1056813>>. [Accessed: 20250428].

¹⁷ Obsérvense las reglas propuestas por Fernández Vivanco y Sanfilippo (2025:51-52).

- Cregan, Anne; Schwitter, Rolf; Meyer, Thomas (2007). Sydney OWL Syntax—towards a controlled natural language syntax for OWL 1.1. In: *Proceedings of OWLED 2007*. <<https://ceur-ws.org/Vol-258/paper36.pdf>>. [Accessed: 20250428].
- Da Cunha, Iria (2022). *Lenguaje claro y tecnología en la Administración*. Granada: Comares. <<https://doi.org/10.55323/edc.2023.52>>. [Accessed: 20250428].
- Da Cunha, Iria (2024). Inteligencia artificial y lenguaje claro en español. In: Retegui, Andrea (ed.). *Lenguaje claro en Iberoamérica. Principios y prácticas*. Buenos Aires: Thomson Reuters, pp. 417–444.
- De Vecchi, Dardo Mario (2018). Company-speak, organisation-speak. In: Humbley, John; Budin, Gerhard; Laurén, Christer (eds.). *Languages for Special Purposes: An International Handbook*. Basel: De Gruyter Mouton, pp. 279–288. <<https://doi.org/10.1515/9783110228014-014>>. [Accessed: 20250428].
- DeepL (2024). About DeepL language models. <<https://support.deepl.com/hc/en-us/articles/14241705319580-About-DeepL-language-models>>. [Accessed: 20250428].
- Feng, Zhaopeng *et al.* (2024). TEaR: Improving LLM-based Machine Translation with Systematic Self-Refinement. *arXiv*. <<https://doi.org/10.48550/arXiv.2402.16379>>. [Accessed: 20251216].
- Fernández Vivanco, Andrea; Bustos Gisbert, José M. (2025). *Lenguajes controlados aplicados a la traducción automática: corpus y evaluación* [Data set]. Zenodo. <<https://doi.org/10.5281/zenodo.17521360>>. [Accessed: 20250428].
- Fernández Vivanco, Andrea; Sanfilippo, Vincenzo (2025). La lecturabilidad de textos escritos en lengua española: revisión crítica de métricas y propuesta de nuevas variables. *ELUA*, v. 44, pp. 29–56. <<https://doi.org/10.14198/ELUA.29266>>. [Accessed: 20250428].
- Funk, Adam *et al.* (2007). CLOnE: Controlled language for ontology editing. In: *Proceedings of ISWC 2007 + ASWC 2007*, pp. 142–155. <https://doi.org/10.1007/978-3-540-76298-0_11>. [Accessed: 20251216].
- García Muñoz, Óscar (2012). *Lectura fácil: métodos de redacción y evaluación*. Madrid: Real Patronato sobre Discapacidad. <<https://www.plenainclusion.org/wp-content/uploads/2021/03/lectura-facil-metodos.pdf>>. [Accessed: 20251216].
- Guilou, Liane; Hardmeier, Christian (2016). PROTEST: A Test Suite for Evaluating Pronouns in Machine Translation. In: *Proceedings of the Tenth International Conference on Language Resources and Evaluation (LREC'16)*, pp. 636–643. <<https://aclanthology.org/L16-1100/>>. [Accessed: 20251216].
- Huijsen, Willem Olaf (1998). Controlled Language – An Introduction. In: *Proceedings of CLAW 98*, pp. 1–15.
- Intento (2025). *The State of Translation Automation 2025*. <<https://inten.to/the-state-of-translation-automation-2025/>>. [Accessed: 20250428].

- Jiménez Hurtado, Catalina; Medina Reguera, Ana (2022). Metodología de la traducción a lectura fácil: Retos de investigación. In: Castillo Bernal, María P. (ed.). *Translation, Mediation and Accessibility for Linguistic Minorities*. Berlin: Frank & Timme, pp. 205–222. <<https://doi.org/10.26530/20.500.12657/54058>>. [Accessed: 20250428].
- Kittredge, Richard (2014). Sublanguages and Controlled Languages. In: Mitkov, Ruslan (ed.). *The Oxford Handbook of Computational Linguistics*. 2nd ed. Oxford: Oxford Academic, pp. 454–472. <<https://doi.org/10.1093/oxfordhb/9780199573691.001.0001>> [Accessed: 20250428].
- Kittredge, Richard; Lehrberger, John (1982). *Sublanguages: Studies of Language in Restricted Semantic Domains*. Berlin / New York: De Gruyter. <<https://doi.org/10.1515/9783110844818>>. [Accessed: 20250428].
- Kuhn, Tobias (2014). A Survey and Classification of Controlled Natural Languages. *Computational Linguistics*, v. 40, n. 1, pp. 121–171. <https://doi.org/10.1162/COLI_a_00168>. [Accessed: 20250428].
- Kuhn, Tobias B.; Barbano, Paolo E.; Nagy, Mate Levente.; Krauthammer, Michael (2013). Broadening the scope of nanopublications. In: *Proceedings of ESWC 2013*, pp. 487–501. <https://doi.org/10.1007/978-3-642-38288-8_33>. [Accessed: 20250428].
- Lohar, Pintu; Xie, Guangyi; Gallagher, Daniel; Way, Andy (2023). Building Neural Machine Translation Systems for Multilingual Participatory Spaces. *Analytics*, v. 2, n. 2, pp. 393–409. <<https://doi.org/10.3390/analytics2020022>>. [Accessed: 20250428].
- Montolío Durán, Estrella (2023). *Guía de redacción judicial clara*. Madrid: Ministerio de la Presidencia, Justicia y Relaciones con las Cortes.
- Núñez Ladevéze, Luis (1993). *Teoría y práctica de la construcción del texto*. Barcelona: Ariel.
- Nyberg, Eric; Mitamura, Teruko; Huijsen, Willem Olaf (2003). Controlled language for authoring and translation. In: Sommers, Harold (ed.). *Computers and Translation: A Translator's Guide*. Amsterdam: John Benjamins, pp. 245–281. <<https://doi.org/10.1075/btl.35.17nyb>>. [Accessed: 20250428].
- O'Brien, Sharon (2003). Controlling controlled English. In: *EAMT Workshop: Improving MT through Other Language Technology Tools*. Budapest: European Association for Machine Translation. <<https://aclanthology.org/2003.eamt-1.12/>>. [Accessed: 20251216].
- O'Brien, Sharon (2006). Eye-tracking and translation memory matches. *Perspectives: Studies in Translation Theory and Practice*, v. 14, n. 3, pp. 185–205. <<https://doi.org/10.1080/09076760708669037>>. [Accessed: 20250428].
- O'Brien, Sharon (2010). Controlled Language and Readability. In: Shreve, Gregory; Angelone, Erik (eds.). *Translation and Cognition*. Amsterdam: John Benjamins, pp. 143–165. <<https://doi.org/10.1075/ata.xv>>. [Accessed: 20251216].
- O'Brien, Sharon (2019). Controlled Language and Writing for an International Audience. In: Maylath, Bruce; St.Amant, Kirk (eds.). *Translation and Localization: A Guide for*

- Technical and Professional Communicators*. New York: Routledge, pp. 95–119. <<https://doi.org/10.1080/10572252.2021.1915064>>. [Accessed: 20251216].
- Pease, Adam; Li, John (2010). Controlled English to logic translation. In: Poli, Roberto; Healy, Michael; Kameas, Achilles (eds.). *Theory and Applications of Ontology: Computer Applications*. Amsterdam: Springer, pp. 245–258. <https://doi.org/10.1007/978-90-481-8847-5_11>. [Accessed: 20250428].
- Pellicer, Lluís (2024). España se enfrenta a otra dana solo dos semanas después de la catástrofe de Valencia. *El País*. <<https://elpais.com/espana/2024-11-12/espana-se-prepara-para-otra-dana-solo-dos-semanas-despues-de-la-catastrofe-de-valencia.html>>. [Accessed: 20240912].
- Pistola Grille, Sara; Viñuales-Ferreiro, Susana (2021). Una clasificación actualizada de los géneros textuales de la Administración pública española. *Revista de Llengua i Dret*, v. 75, pp. 181–203. <<https://doi.org/10.2436/rld.i75.2021.3587>>. [Accessed: 20250428].
- Pratt-Hartmann, Ian (2003). A two-variable fragment of English. *Journal of Logic, Language and Information*, pp. 13–45. <<https://doi.org/10.1023/A:1021149027971>>. [Accessed: 20250428].
- Ramírez Polo, Laura (2012a). Los lenguajes controlados y la documentación técnica: mejorando la traducibilidad. *Tradumàtica*, v. 10, pp. 192–204. <<https://doi.org/10.5565/rev/tradumatica.25>>. [Accessed: 20250428].
- Ramírez Polo, Laura (2012b). *Use and Evaluation of Controlled Languages in Industrial Environments and Feasibility Study for the Implementation of Machine Translation* [Tesis doctoral]. Universitat de València.
- Real Academia Española (2024). *Guía panhispánica de lenguaje claro y accesible*. Madrid: Espasa.
- Roturier, Johann (2006). *An Investigation into the Impact of Controlled English Rules on the Comprehensibility, Usefulness and Acceptability of Machine-Translated Technical Documentation for French and German Users* [Tesis doctoral]. Dublin City University.
- Ruiz Cascales, Remedios; Sutcliffe, Richard F. (2003). A specification and validating parser for simplified technical Spanish. In: *EAMT Workshop: Improving MT through Other Language Technology Tools*. Budapest: European Association for Machine Translation. <<https://aclanthology.org/2003.eamt-1.15/>>. [Accessed: 20251216].
- Saggion, Horacio *et al.* (2015). Making It Simplex: Implementation and Evaluation of a Text Simplification System for Spanish. *ACM Transactions on Accessible Computing*, v. 6, n. 7, pp. 1–36. <<https://doi.org/10.1145/2738046>>. [Accessed: 20250428].
- Sánchez-Gijón, Pilar; Kenny, Dorothy (2022). Selecting and preparing texts for machine translation: Pre-editing and writing for a global audience. In: Kenny, Dorothy (ed.). *Machine Translation for Everyone*. Berlin: Language Science Press, pp. 81–103.

- Schieferdecker, Ina (2024). Next-Gen Software Engineering. Big Models for AI-Augmented Model-Driven Software Engineering. *arXiv*. <<https://doi.org/10.48550/arXiv.2409.18048>>. [Accessed: 20250428].
- Schwiter, Rolf (2002). English as a formal specification language. In: *Proceedings of DEXA '02*, pp. 228–232. <<https://doi.org/10.1109/DEXA.2002.1045903>>. [Accessed: 20250428].
- Seibel, Claudia; Carlucci, Laura (2021). El lenguaje controlado: punto de partida hacia la Lectura fácil. *HIKMA*, v. 20, n. 2, pp. 331–357. <<https://doi.org/10.21071/hikma.v20i2.13394>>. [Accessed: 20250428].
- Sowa, John (2000). Ontology, metadata, and semiotics. In: *Proceedings of ICCS 2000*, pp. 55–81. <https://doi.org/10.1007/10722280_5>. [Accessed: 20250428].
- Temnikova, Irina (2012). *Text Complexity and Text Simplification* [Tesis doctoral]. University of Wolverhampton.
- Williams, Christopher (2023). *The Impact of Plain Language on Legal English in the United Kingdom*. New York: Routledge.
- Wu, Yonghui *et al.* (2016). Google's Neural Machine Translation System: Bridging the Gap between Human and Machine Translation. *arXiv*. <<https://doi.org/10.48550/arXiv.1609.08144>>. [Accessed: 20250428].