

ERRORES LÉXICOS EN EL ESPAÑOL ORAL NO NATIVO: ANÁLISIS DE LA INTERLENGUA BASADO EN CORPUS

LEONARDO CAMPILLOS LLANOS
Universidad Autónoma de Madrid
leonardo.campillos@uam.es

Resumen

Se analizan los errores léxico-semánticos en la producción oral de cuarenta aprendices de español. Los datos pertenecen a un corpus de interlengua compuesto por entrevistas con universitarios de más de 9 lenguas maternas y de nivel intermedio (A2 y B1, Marco Común Europeo de Referencia). La metodología es la investigación en corpus de estudiantes, en concreto el análisis de errores asistido por ordenador. Los resultados muestran que los errores formales son más frecuentes en A2, y que no abundan los semánticos, pero persisten aumentando ligeramente en B1.

PALABRAS CLAVE: errores léxico-semánticos; investigación en corpus de estudiantes; español como lengua extranjera; interlengua oral.

Abstract

This study analyses the lexical errors in the oral production of forty learners of Spanish. Data belong to a learner corpus of oral interviews with university learners from over nine language backgrounds at intermediate level: A2 (N=20) and B1 (N=20) (Common European Framework of Reference). Our methodology is Learner Corpus Research, specifically Computer-aided Error Analysis. The results in our corpus show that formal errors are more frequent at A2 level, whereas semantic errors, although being less abundant, persist and slightly increase at B1.

KEY WORDS: Lexical-semantic errors; Learner corpus research; Spanish as a Foreign Language; Oral interlanguage.

1. Introducción*

La adquisición del léxico en una lengua segunda/extranjera suscita un enorme interés tanto en profesores como alumnos. No solamente respecto a cuestiones como la cifra de palabras más adecuada que formaría el vocabulario de cada nivel, o las técnicas más efectivas para acelerar su aprendizaje, sino también por la curiosidad de comprender la creatividad léxica de cada aprendiz y su capacidad de *lograr* expresar el mensaje. Muchas veces el empeño de la didáctica se centra *lo gramatical*, pero lo cierto es que, aunque los errores de gramática sean más abundantes, realmente suelen ser los errores de vocabulario los que impiden la comunicación —o, cuando menos, desdibujan los matices del enunciado—.

* Trabajo financiado por la Comunidad de Madrid y el Fondo Social Europeo mediante un contrato de investigación predoctoral. Agradezco sinceramente a los revisores anónimos las sugerencias para la mejora del artículo.

Por lo que concierne al español, las investigaciones son abundantes en este aspecto. Sin ánimo de exhaustividad, destacamos trabajos de distinta naturaleza: Lafford y Collentine (1987), Baralo (1996; 2001), Fernández Silva (1998), Collentine (2004), Whitley (2004) o Marsden y David (2008). En el nivel léxico también se han adentrado los estudios acerca de las *estrategias de comunicación oral*, dado que son mecanismos de los que se vale el alumno ante un déficit léxico (Tarone, 1980; Poullisse, 1990; y para el español, Liang y Llobera, 1996; Pinilla, 2000; Muñoz Benito, 2002; Lafford, 2004; y Rubio, 2005).

Este artículo pretende contribuir al conocimiento del proceso de adquisición del léxico en el español oral por alumnos de nivel intermedio-bajo, aunque desde el enfoque del Análisis de errores (Corder, 1971), y siguiendo la metodología de Fernández López (1990a, 1990b). Se trata de un resumen del análisis descrito en Campillos Llanos (2012), al que remitimos para ampliar los puntos que se exponen.

2. Presentación de la investigación

Este estudio se enmarca en la investigación con corpus de estudiantes (*Learner Corpus Research*; Granger, 2012), y específicamente, el análisis de errores asistido por ordenador (*Computer-aided Error Analysis*; Dagneaux *et alii*, 1998; Barlow, 2005). Dicha metodología se fundamenta en el análisis de errores (AE) (Corder, 1971): identificación, descripción, clasificación, diagnóstico y valoración del error. Es cierto que esta metodología tiene limitaciones para abordar la *evitación* del error por los no nativos, o para distinguir *fallos* (debidos a la actuación) y *errores* (debidos a la competencia; Fernández López, 1990a; Granger, 2003). Con todo, el AE es óptimo para abordar la causa del error, y el etiquetado informático permite recuperar errores según criterios de búsqueda y extraer las frecuencias de cada tipo.

Asimismo, usamos la metodología de corpus, que se ha venido aplicando a la didáctica e investigación lingüística (Cruz Piñol, 2012). Para el español, existe el corpus de textos CEDEL2 (Lozano, 2009), y respecto a la oralidad, contamos con el corpus longitudinal de entrevistas recogidas por Díaz Rodríguez (2007) y los corpus Spanish Learner Language Oral Corpora (SPLLOC; Mitchell *et alii*, 2008).

2.1. Participantes y diseño del corpus

Los participantes fueron cuarenta (N = 40) universitarios extranjeros que estudiaban español en Madrid mediante convenios internacionales (p. ej., ERASMUS) y otros programas de intercambio. Todos (salvo una alumna) tenían entre 19 y 26 años, y cursaban grado o posgrado¹. El marco de investigación, pues, se restringe al uso del español en contexto académico. Por lo que respecta al nivel, la mitad de los estudiantes

¹ Algunos alumnos habían recibido instrucción formal en sus países, y otros habían comenzado en Madrid, combinando instrucción formal y adquisición en contexto natural. Todos habían estudiado en un contexto de inmersión varios meses antes de la entrevista.

(N = 20) cursaba el nivel plataforma (A2), y la otra mitad, el nivel umbral (B1) (*Marco Común Europeo de Referencia*, en adelante, MCER; Consejo de Europa, 2001). Se entrevistó a cuatro sujetos por cada grupo de lengua materna (en adelante, L1), reuniendo datos de más de 9 orígenes lingüísticos de tipología variada: lenguas románicas (italiano, francés y portugués), germánicas (inglés, alemán y neerlandés), eslavas (polaco), sino-tibetanas (chino) y japonés. Hay otro grupo mixto de cuatro alumnos con otras L1 (finés, coreano, turco y húngaro). Debido a la disponibilidad de los alumnos, no se alcanzó en los 10 grupos un equilibrio en cuanto al nivel o al sexo del hablante.

Cada entrevista duró en torno a 15-20 minutos, y se grabó más de una hora para cada grupo (en total, más de trece horas). Las grabaciones se recogieron entre los cursos 2007-08, 2008-09 y 2009-2010.

2.2. Método de obtención de datos

Los informantes fueron grabados durante una entrevista semiestructurada con el investigador. La participación era voluntaria (con consentimiento firmado), y después recibían una explicación sobre sus errores. Tras una presentación (en la que los alumnos informaban acerca de sus estudios, conocimiento de otras lenguas o tiempo de aprendizaje de español), todos realizaban las mismas tareas, para obtener datos comparables. Las técnicas de obtención de muestras fueron similares a las de los exámenes de idiomas:

- Descripción de dos fotografías (dos tipos de comida).
- Una tarea narrativa a partir de dos historias representadas en viñetas, a fin de que los alumnos usaran tiempos de pasado y mecanismos de cohesión discursiva, y para plantearles una pregunta sobre dos actos de habla (una sugerencia y una petición).

Al final de la entrevista, los participantes tenían que expresar su opinión acerca de diferentes asuntos (p. ej., los cambios en los hábitos de alimentación actual), con el fin de obtener un habla más espontánea.

2.3. Procesamiento y análisis de datos

Las entrevistas fueron transcritas manualmente mediante el programa Transana² y siguiendo las convenciones de los proyectos C-Oral-Rom (Cresti y Moneglia, 2005) y SPLLOC. Las transcripciones se anotaron mediante GRAMPAL, un analizador morfológico para el español adaptado para el tratamiento de datos orales (Moreno Sandoval y Guirao, 2006). GRAMPAL procesa locuciones (p. ej. *gracias a*) así como marcadores discursivos (ej. *quiero decir*). Se llevó a cabo una revisión manual de la anotación a fin de corregir ambigüedades debidas a la asignación automática de categorías. La anotación morfológica permitió extraer listados de frecuencia de uso de categorías y de unidades léxicas.

² Chris Fassnacht y David Woods (University of Wisconsin). <http://www.transana.org>

Solo se obtuvieron recuentos de los turnos de los alumnos. La normalización de recuentos de errores utiliza como base la *unidad léxica*, considerando como tal la palabra o conjunto de palabras con unidad de significado. Por ejemplo, *escuela de idiomas* cuenta como tres palabras, pero como una unidad léxica. Consideramos una sola unidad léxica los tiempos compuestos y las perífrasis (ej. *hubiera venido, suele venir*), las locuciones (ej. *echar de menos*), los marcadores discursivos (ej. *vamos a ver*), las fórmulas oracionales (ej. *y ya está*) y los nombres propios (ej. *Puerto Rico*). Los alumnos produjeron 52688 unidades léxicas, con una distribución por niveles prácticamente semejante: 26317 en A2, y 26371 en B1.

Paralelamente a la anotación morfológica, se realizó el análisis y el etiquetado de errores léxicos. Para ello, se diseñó una tipología de errores a partir de trabajos previos (Fernández López, 1990a; Vázquez, 1999; *vid.* resumen en Díaz-Negrillo y Fernández, 2006) y se confeccionó al efecto un conjunto de etiquetas en lenguaje XML (W3C, 2012). Los errores fueron etiquetados manualmente y los documentos se integraron en una interfaz informática para la consulta y el recuento de cada tipo de error³

3. Errores léxicos

El recuento de errores no ambiguos que se consideró de tipo léxico es de 2008 (29,37% del total de 6838 errores), a los cuales hay que añadir los 244 ambiguos que pueden juzgarse en dos niveles (principalmente, el fonético, aunque también el gramatical). Se trata, pues, de errores que se registraron con alta frecuencia en el corpus, y que en gran medida dificultaron comprender el mensaje. Si se analiza la cifra de errores léxicos por niveles, se constata que decrecieron del nivel básico al intermedio. Un test U de Mann-Whitney bilateral (los datos no tenían una distribución normal) confirmó que el descenso de errores no ambiguos fue estadísticamente significativo ($p = 0.046$, < 0.05). Así, en A2 se registraron 1309 errores más 177 ambiguos. Los 1309 errores no ambiguos representan un 32,37% sobre el total de 4044 incorrecciones de A2, y afectaron al 4,97% de las unidades léxicas (*uu. ll.* en adelante) de dicho nivel (26317). En cambio, en B1 se registraron 699 errores, junto a 67 ambiguos. Los 699 errores no ambiguos equivalen a un total del 25,02% de todos los errores de B1 (2794), y afectaron al 2,65% de las *uu. ll.* de dicho nivel (26371). Los datos se representan en las tablas 1 y 2, y en los gráficos 1 y 2 (en la tabla 1, y en adelante, *M* es la abreviatura de *Media*, y *D* es la abreviatura de *desviación estándar*).

³ El acceso a la interfaz se encuentra disponible en: <http://cartago.lllf.uam.es/corele/index.html>. En esta dirección se pueden consultar las entrevistas completas para comprender adecuadamente (en su contexto de aparición) los ejemplos que presentamos.

Nivel	ERRORES (VALORES ABSOLUTOS)					
	No ambiguos			Ambiguos		
	Errores	M	D	Errores	M	D
A2	1309	65,45	70,86	177	8,85	14,36
B1	699	34,95	28,14	67	3,35	4,07
Total	2008	50,20	55,41	244	6,10	10,79

Tabla 1. Frecuencia de errores léxicos (no ambiguos y ambiguos) según el nivel

Nivel	ERRORES (PORCENTAJES)			
	No ambiguos		Ambiguos	
	% sobre total de errores	% sobre uu. ll.	% sobre total de errores	% sobre uu. ll.
A2	32,37% (sobre 4044)	4,97% (26317)	4,38% (sobre 4044)	0,67% (26317)
B1	25,02% (sobre 2794)	2,65% (26371)	2,40% (sobre 2794)	0,25% (26371)
Total	29,37% (sobre 6838)	3,81% (52688)	3,57% (sobre 6838)	0,46% (52688)

Tabla 2. Frecuencia relativa de errores léxicos y proporción sobre unidades léxicas de cada nivel

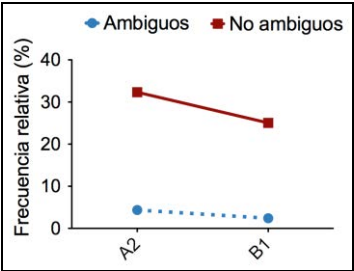


Figura 1. Frecuencia relativa de errores léxicos por nivel (% sobre el total de errores)

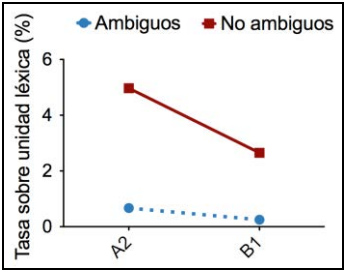


Figura 2. Tasa de errores léxicos por unidades léxicas de cada nivel (% sobre uu. ll. de cada nivel)

La tabla 3 muestra los recuentos de error por grupos de L1. Los valores entre paréntesis de las columnas de frecuencia relativa muestran el total de errores de cada grupo. Los valores entre paréntesis de las columnas «% sobre uu.ll.» corresponden a la cifra de unidades léxicas de cada grupo. Sobresale el número de errores del grupo lusófono, con una media (en adelante, *M*) de 178,75 errores no ambiguos por entrevista, seguido del germanohablante (*M* = 51,25) y el neerlandés (*M* = 51,25). No obstante, los valores de la desviación estándar (en adelante, *D*) son muy altos en estos grupos: respectivamente, *D* = 87,71, *D* = 67,25, y *D* = 41,12 (en todos los casos nos referimos solo a los errores no ambiguos). Esto quiere decir que ciertos alumnos presentaron altas tasas de error en comparación con otros del mismo grupo de L1 (alta variación intragrupal). El grupo chino no se queda atrás en número de errores (*M* = 33,25, *D* = 5,38), especialmente si se consideran también los ambiguos (figura 3). Por el contrario, produjeron menos errores el grupo italiano (*M* = 27,50, *D* = 7,94) y el grupo heterogéneo (*M* = 25,50, *D* = 9,95).

No ambiguos						Ambiguos				
L1	Errores	M	D	Frec. relativa (% del total de errores del grupo)	% sobre uu. ll. del grupo	Errores	M	D	Frec. relativa (% del total de errores del grupo)	% sobre uu. ll. del grupo
Portugués	715	178,7 5	87,71	54,41% (de 1314)	9,76% (de 7330)	90	22,50	28,49	6,85% (de 1314)	1,23% (de 7330)
Alemán	205	51,25	67,25	32,08% (de 639)	4,48% (de 4573)	9	2,25	0,96	1,41% (de 639)	0,20% (de 4573)
Neerlandés	205	51,25	41,12	28,16% (de 728)	3,95% (de 5186)	38	9,50	5,26	5,22% (de 728)	0,73% (de 5186)
Chino	133	33,25	5,38	15,56% (de 855)	3,70% (de 3592)	32	8,00	5,35	3,74% (de 855)	0,89% (de 3592)
Inglés	136	34,00	4,32	24,42% (de 557)	2,87% (de 4731)	45	11,25	6,95	8,08% (de 557)	0,95% (de 4731)
Japonés	130	32,50	2,38	19,97% (de 651)	2,79% (de 4664)	12	3,00	1,83	1,84% (de 651)	0,26% (de 4664)
Polaco	124	31,00	10,36	20,43% (de 607)	2,34% (de 5300)	1	0,25	0,50	0,16% (de 607)	0,02% (de 5300)
Francés	148	37,00	16,67	29,42% (de 503)	2,31% (de 6413)	4	1,00	1,15	0,80% (de 503)	0,06% (de 6413)
Otros	102	25,50	9,95	16,37% (de 623)	2,14% (de 4767)	8	2,00	1,83	1,28% (de 623)	0,17% (de 4767)
Italiano	110	27,50	7,94	30,47% (de 361)	1,79% (de 6133)	5	1,25	2,50	1,39% (de 361)	0,08% (de 6133)
Total	2008	50,20	55,41	29,37% (de 6838)	3,81% (de 52688)	244	6,10	10,79	3,57% (de 6838)	0,46% (de 52688)

Tabla 3. Distribución de errores léxicos por grupos de lengua materna

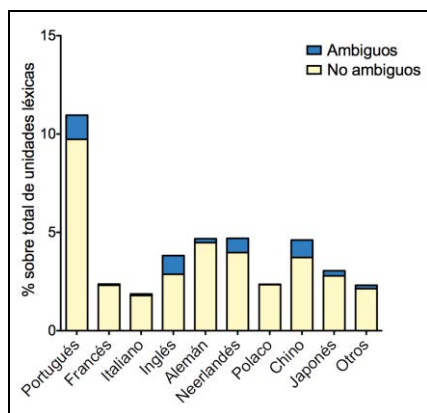


Figura 3. Errores léxicos por grupos de L1 (% sobre uu. ll. de cada grupo)

Las razones de estos datos pueden ser de varios tipos, aunque no están del todo claras. En primer lugar, debe considerarse el nivel lingüístico de los estudiantes, ya que tres de los cuatro participantes pertenecían al nivel A2 en el grupo neerlandés y portugués (aunque no en el alemán). Como consecuencia, el número de errores aumenta proporcionalmente, y esto puede explicar la alta variación intragrupal. Por otra parte, podría deberse a la distancia lingüística entre el español y las lenguas germánicas, cuyo léxico posee raíces de distinto origen. Dicho lo cual, esto tampoco explica por qué se registraron más errores en el grupo alemán (perteneciente, como el español, a la familia indoeuropea) que en el grupo chino (lengua de la familia sinotibetana, y más alejada tipológicamente). Tampoco las razones tipológicas explican por qué el grupo portugués cometió con diferencia más errores que el grupo italiano o francés, pese a ser todas estas lenguas románicas como el español. La interferencia parece influir muy negativamente entre los aprendices lusófonos, pero, como señalamos, los valores de la desviación son muy altos. Al no contar con suficientes informantes de cada lengua materna, apenas pueden esbozarse conclusiones sólidas solo con las dos variables del nivel y la lengua materna. Además, creemos que tienen gran influencia variables de tipo individual (por ejemplo, habría que analizar el tiempo de estudio de español o de estancia en un país hispanohablante de los estudiantes en cada grupo).

Resulta interesante observar los tipos de error. Respecto a los no ambiguos, registramos 1621 de forma ($M = 40,52$, 80.73% de errores léxicos) y 384 de significado ($M = 9,60$, 19,12% del total de errores léxicos) (tabla 4). Junto a estos no incluimos 3 incorrecciones léxicas que podrían considerarse de forma o significado. En cuanto a los errores ambiguos (esto es, que también podrían relacionarse con la gramática o la pronunciación), recogimos 244 casos que pueden ser tanto de forma como de significado. Aunque la clasificación de errores en formales o semánticos se basa principalmente en la tipología de Fernández (1990a), los resultados que obtuvimos en la producción oral difieren de los textos escritos. En nuestros datos orales, la mayor parte de incorrecciones se deben a la forma y no al significado.

La distribución de errores de forma y de significado en nuestros datos orales muestra que predominan las incorrecciones formales en A2 (1124 errores no ambiguos, $M = 56,20$) (tabla 4). Dichas dificultades van desapareciendo en B1, reduciéndose a más de la mitad (497 errores no ambiguos, $M = 24,85$). Por el contrario, los errores de significado, pese a que se registran en menor número en ambos niveles, en nuestro corpus se mantienen (aumentando ligeramente) al contrastar el nivel A2 con el B1 (respectivamente, 182 errores, $M = 9,10$, y 202 errores, $M = 10,10$). El gráfico 4 muestra la distribución de los errores de forma y significado (en porcentaje sobre uu. ll. de cada nivel: 26317 de A2 y 26371 de B1).

Nivel		No ambiguos			Total	Ambiguos
		Forma	Significado	No asignado		
A2	N.º errores	1124	182	3	1309	177
	M	56,20	9,10	0,15	65,45	8,85
	% sobre uu. ll.	4,27%	0,69%	0,01%	4,97%	0,67%
B1	N.º errores	497	202	0	699	67
	M	24,85	10,10	0,00	34,95	3,35
	% sobre uu. ll.	1,88%	0,77%	0,00%	2,65%	0,25%
Total	N.º errores	1621	384	3	2008	244
	M	40,52	9,60	0,07	50,20	6,10
	% sobre uu. ll.	3,08%	0,73%	0,01%	3,81%	0,46%
	% total errores léxico	80,73%	19,12%	0,15%	-	-

Tabla 4. Distribución de errores léxicos por niveles (forma, significado y otros)

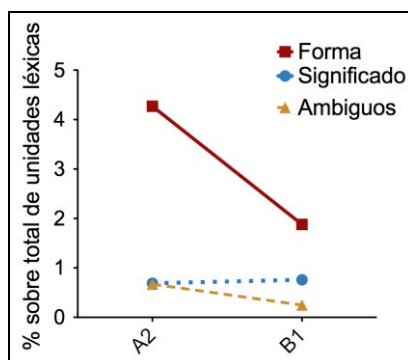


Figura 4. Tipos de errores léxicos (forma y significado) por nivel

3.1. Errores léxicos de forma

En conjunto, los errores formales afectaron al 3,08% del total de unidades léxicas producidas (52688). La tabla 5 y el gráfico 5 resumen la clasificación de los errores de forma. Estas cifras han de tomarse con cautela, ya que, en determinados casos, el tipo de

error se puede prestar a diferentes interpretaciones. Los errores más abundantes fueron los extranjerismos: 875 errores no ambiguos (M = 21,87, D = 46,99, un 43,58% sobre el total de errores léxicos). No obstante, el valor de la desviación es muy alto debido a que dichos errores se concentraron en los grupos lusófono, germano y neerlandés. Aparecen en segundo lugar en cuanto a frecuencia las formaciones no atestiguadas: 359 no ambiguos (M = 8,97, D = 7,64, 17,88% sobre el total de errores léxicos). Y a continuación, los errores de asignación de género: 272 no ambiguos (M = 6,80, D = 5,54, 13,55% del total de errores léxicos).

ERRORES LÉXICOS DE FORMA (NO AMBIGUOS)				
	Errores	M	D	% sobre total de errores léxicos (2008)
Extranjerismos	875	21,87	46,99	43,58%
Formaciones no atestiguadas	359	8,97	7,64	17,88%
Género	272	6,80	5,54	13,55%
Número	53	1,32	1,67	2,64%
Malapropismos	23	0,58	0,84	1,15%
Otros	23	0,58	1,78	1,15%
Calcos	16	0,40	0,84	0,80%

Tabla 5. Distribución de errores léxicos de forma no ambiguos

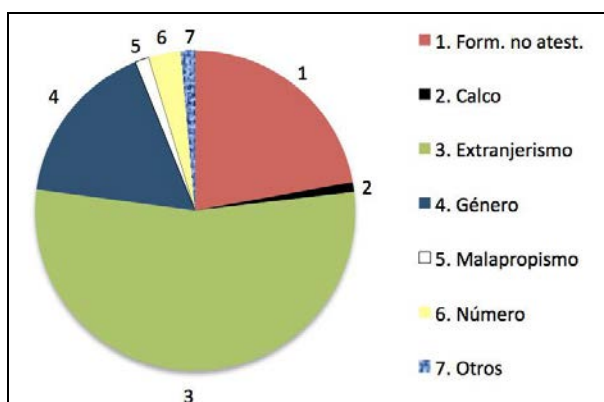


Figura 5. Distribución de errores léxicos de forma no ambiguos (% sobre total de errores de forma)

3.1.1. Malapropismo

Recogemos aquí las confusiones entre palabras con una similitud fonética o formal (parónimos), y que se han denominado *uso de un signifiicante próximo* en otros trabajos (Fernández López, 1990a). Según Fat y Cutler (1977; *vid.* Hernández Fernández, 2004), son caracterizados por tres rasgos: primero, la palabra errónea existe en la lengua; segundo, sus significados no están relacionados; y tercero, poseen una pronunciación muy parecida (a veces incluso el mismo número de sílabas). Los errores por

malapropismo también los cometen los propios nativos (son fallos mecánicos o de actuación), y reafirman la hipótesis de la organización fonológica del lexicón mental (Fat y Cutler, 1977).

La clasificación de algunos malapropismos no es fácil. Por ejemplo, la confusión entre la forma *di* (imperativo de *decir* o pretérito de *dar*) y *dije* o *dijo* (pretérito de *decir*) no es claro si es un error de conjugación o de malapropismo. Otros casos se podrían considerar un acortamiento de palabra (p. ej., *seguir* por *conseguir*).

Según nuestro criterio, se produjeron 23 malapropismos no ambiguos (M = 0,58, D = 0,84), que afectan apenas a un 0,04% de las unidades léxicas totales. Hubo también 2 casos ambiguos. Su distribución por niveles es casi exactamente igual: en A2 se registraron 12 no ambiguos y 2 ambiguos, y en B1, 11 casos no ambiguos. Dicha distribución puede reflejar que se trata de un error poco relacionado con el nivel de competencia (véanse los ejemplos; indicamos el nombre del documento entre paréntesis).

- (1) ADC: y les &eh &mm / **convertimos** que [/] que podemos hhh {%act: laugh} dormir aquí ('convencimos') (POLMB1)
- (2) FAN: dos / persona / &mm [/] &ah / un [/] un / &mm [/] &pa [/] &pa [/] pájaro [/] &paj [/] ¿pajero ? // **pájara** ... ('pareja') (CHIWA2_2)

También hay dos casos ambiguos; por ejemplo, el siguiente contexto puede ser una por confusión con el verbo *vivir*, pero también una pronunciación de *e* muy cerrada ([i]):

- (3) STE: &eh quiere → **beber** {%pho: [bi'bir]} vino /// (DUTMA2)

Podemos concluir comentando que el diagnóstico de estos errores no es interlingüístico, pues la confusión aparece entre significantes de la propia lengua meta. Ahora bien, no queda claro si se trata de un problema de la competencia lingüística del alumno, o un lapsus puntual de acceso al lexicón. A pesar de que no parece presentar gran frecuencia de aparición, en ciertos contextos podría repercutir negativamente en la comunicación.

3.1.2. Asignación de género erróneo

Siguiendo a Fernández López (1990a), consideramos el género como rasgo léxico del nombre, así que no incluimos en este apartado las concordancias erróneas (ej. *la plaza *nuevo*). En ocasiones no es claro si la incorrección se debe al género o a la concordancia. Como criterio, marcamos aquí los casos en que el determinante (ej., **la gesto*), el adjetivo (ej., *filete *picada*) o la vocal *-a/-o* (ej., **dinera*) expresan el uso incorrecto de masculino o femenino. Tampoco agrupamos aquí otros errores por formación del género gramatical (p. ej., **profesoro*), ni los casos que implican un cambio de significado, que abordamos en el apartado de relación semántica (p. ej., *bolso* frente a *bolsa*).

Según estos criterios, registramos 272 errores no ambiguos (M = 6,8, D = 5,54, 13,55% de errores léxicos). La distribución en niveles es muy similar: hubo más errores en A2 (148, M = 7,4, D = 5,67, el 54,41% de los errores de género), y 124 en B1 (M = 6,2, D = 5,49, el 45,59%) (tabla 6 y figura 6).

Nivel	No ambiguos					Ambiguos				
	Errores	Frec. rel. %	M	D	% sobre uu. ll.	Errores	Frec. rel. %	M	D	% sobre uu. ll.
A2	148	54,41%	7,4	5,67	0,56%	6	75%	0,3	0,66	0,023%
B1	124	45,59%	6,2	5,49	0,47%	2	25%	0,1	0,31	0,008%
Total	272	100,00%	6,8	5,54	0,52%	8	100%	0,2	0,52	0,015%

Tabla 6. Distribución por niveles de errores de género no ambiguos y ambiguos

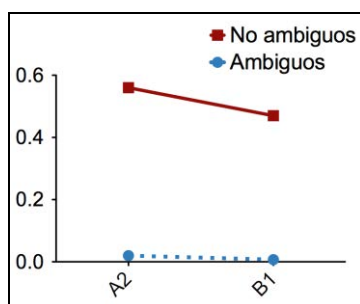


Figura 6. Errores de género por nivel (% sobre uu. ll. de cada nivel)

Por grupos, se recogieron estas cifras de errores no ambiguos (se ordenan por la tasa sobre uu. ll.):

L1	Errores	M	D	% sobre total de uu. ll. del grupo	Un error cada...
Inglés	39	9,75	3,59	0,82%	121,30 unidades.
Japonés	38	9,50	7,33	0,81%	122,73 "
Polaco	43	10,75	5,74	0,81%	123,25 "
Neerlandés	39	9,75	8,81	0,75%	132,97 "
Alemán	31	7,75	5,91	0,68%	147,51 "
Chino	20	5,00	3,56	0,56%	179,60 "
Otros	26	6,50	2,38	0,55%	183,34 "
Francés	24	6,00	4,97	0,37%	267,20 "
Portugués	11	2,75	2,22	0,15%	666,27 "
Italiano	1	0,25	0,50	0,02%	6133,00 "

El grupo italiano, lusófono y francófono son los que produjeron menos errores, mientras que los grupos anglófono, polaco y japonés son los que presentaron una tasa más alta de incorrecciones por unidad léxica. La ausencia de expresión de género en inglés y japonés puede explicar esta tasa tan elevada. Sin embargo, en polaco también

existe el género léxico, y los alumnos con esta L1 cometieron numerosos errores, mientras que en el grupo chino no se registraron tantos, y tampoco se expresa el género léxico. A continuación se transcriben algunos ejemplos de errores no ambiguos:

- (4) REM: los españoles no hacen como **muchas errores** /// (FREMB1)
- (5) FIN: aprender / un [/] **una idioma** (ENGMB1)
- (6) ADC: empezamos **nuestra viaje** {%alt: vaje} [/] **viaje** de Madrid /// (POLMB1)

Algunas dificultades se presentaron con sustantivos en *-e*, cuyo género es arbitrario en castellano (**nuestra viaje*), aunque las excepciones de la regla también suelen ser causa de error (**una idioma*). Por otra parte, los nombres terminados en consonante ofrecen obstáculos respecto al género (**unos veces*), e igualmente, ciertos nombres propios (**una Burger King*) y conceptos extranjeros (**wasabi fresca*).

La causa de ciertos errores es la interferencia en palabras cuyo género difiere en español y en la L1 del aprendiz (*heterogénicos*; Benedetti, 2002: 155): p. ej., **las errores*, femenino en francés. No obstante, la causa interlingüística no siempre es clara, sobre todo respecto a las palabras heterogénicas sin una correspondencia formal directa entre el español y otra lengua: p. ej., **los frases* (¿del alemán *der Satz*, masculino?). Por otra parte, en los casos en que en la lengua materna del hablante no existe la distinción de género (chino, coreano, finés, húngaro, inglés, japonés o turco), o solo se distingue el neutro (el neerlandés), el uso incorrecto del femenino es quizá un error intralingüístico por falsa hipótesis (p. ej., *baile *clásica*). También puede generalizarse a ciertas excepciones la regla de asignación del masculino o del femenino a sustantivos acabados en *-o* (p. ej., **los manos*) o en *-a* (p. ej., los de origen griego: **la problema*). Sin embargo, en la mayor parte de las ocasiones el género seleccionado es el no marcado (el masculino), tal como comprobó McCarthy (2008). Dicho fenómeno puede obedecer a una simplificación de la regla, a la hipercorrección (p. ej., **el tortilla*), o a una posible interferencia de la L1 (p. ej., en casos en que existe género neutro en la L1, como en neerlandés o alemán, y se asigne el masculino: **un foto*, en alemán *das Bild*). En estos contextos creemos que es demasiado aventurado explicar la causa del error.

También se registraron 8 errores ambiguos (6 en A2, y 2 en B1). Algunos pueden deberse a incorrecciones de pronunciación que afectan a la vocal *-a* del femenino. Aparecieron en alumnos de grupos de lenguas que presentan problemas para la pronunciación clara de vocales átonas (ejemplo 7):

- (7) MAR: es más típico de **las** {%pho: [ləs]} **montañas** ... (DUTWA2_2)

Ninguno de los errores de asignación de género reviste importancia para la comprensión del mensaje, excepto en las confusiones que conllevan un cambio de significado. De dichos errores se registraron 3 casos (ej. *bolso/bolsa*, *pimiento/pimienta*) (apdo. 3.2.1). Con todo, parece que se trata de un aspecto cuya adquisición es más lenta que otros puntos, tanto por factores interlingüísticos como intralingüísticos.

3.1.3. Asignación de número erróneo

Consideramos también rasgo léxico el número del sustantivo, e incluimos los errores que aparecen tanto en nombres propios (*la *Falla*, ‘las Fallas’) como en expresiones o locuciones que el uso ha fijado en singular o plural (*en *condición*, ‘en condiciones’). Algunos errores pueden cambiar el significado de la palabra (*la pasta ~ las pastas*) o afectan al rasgo incontable de sustantivos abstractos (**muchos trabajos* por *mucho trabajo*), de clase o materia (**natas*) o colectivos (**las gentes*). No incluimos casos de formación incorrecta del plural y otros errores relacionados con la pluralidad de los nombres contables.

Según dichos criterios, observamos 53 errores no ambiguos (M = 1,325, D = 1,67), que afectan al 0,10% de las unidades léxicas. Esto es, apareció un error aproximadamente cada 994 unidades. Su distribución en niveles muestra que se registraron casi en igual número: 27 en A2 (M = 1,350, D = 1,73), y 26 en B1 (M = 1,300, D = 1,66) (tabla 7 y gráfico 7). Aunque suponen una cantidad inferior que los de género, parece que ciertas dificultades no desaparecen al pasar de un nivel a otro.

Nivel	No ambiguos				
	Errores	Frec. rel. %	M	D	% sobre uu. ll.
A2	27	50,94%	1,350	1,73	0,103%
B1	26	49,06%	1,300	1,66	0,099%
Total	53	100,00%	1,325	1,67	0,101%

Tabla 7. Distribución por niveles de errores de número no ambiguos

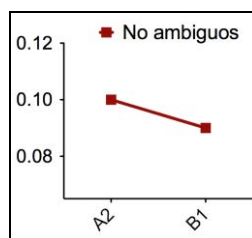


Figura 7. Errores de número por nivel (% sobre uu. ll. de cada nivel)

El número de errores por grupo de L1 es el siguiente (ordenamos según la tasa por uu. ll.):

L1	Errores	M	D	% sobre total de uu. ll. del grupo	Un error cada...
Chino	7	1,75	1,50	0,19%	513,14 unidades.
Neerlandés	10	2,50	2,38	0,19%	518,60 "
Japonés	9	2,25	2,06	0,19%	518,22 "
Inglés	6	1,50	3,00	0,13%	788,50 "
Francés	6	1,50	1,29	0,09%	1068,00 "
Portugués	6	1,50	1,73	0,08%	1221,50 "
Otros	4	1,00	1,41	0,08%	1191,75 "
Italiano	4	1,00	0,82	0,07%	1533,25 "
Polaco	1	0,25	0,50	0,02%	5300,00 "

Transcribimos abajo algunos errores (para más contexto, léanse las transcripciones):

- (8) [El alumno está hablando de los alimentos que compra en el supermercado]
 EME: solamente yo **hago / compras** // y a veces ¡ah! ¿ qué lleva este ?
 (TURWB1)
- (9) NUN: bacalao con **natas** ... (PORMA2)
- (10) AIS: **las gente** aquí hablan inglés /// (ENGWA2)
- (11) AMA: vimos espectáculo de flamenco → &eh **la** <sevillana e todo muy lindo>...
 (PORWB1)

La tasa de error más alta se registró, por un lado, entre los estudiantes chinos y japoneses, y por otro, entre los neerlandeses, aunque por diferentes motivos. Los aprendices asiáticos presentan problemas con la asignación de número porque en sus lenguas maternas no suele expresarse este rasgo, o porque les cuesta comprender la distinción entre sustantivos individuales y colectivos. En cambio, los otros estudiantes producen muchos de estos errores por interferencia de su L1. Así, el error más problemático (por frecuencia y porque se distribuye en grupos de aprendices con diferente L1) aparece en el sustantivo colectivo *gente*, que los anglófonos, francófonos u holandeses concuerdan en plural por transferencia de su L1. Otros errores que merecen atención se registraron en el uso en plural de sustantivos que se usan únicamente en singular en castellano; p. ej., *la compra*, referido al ‘conjunto de alimentos comprados para el consumo doméstico’. Aunque dichos errores podrían causar alguna ambigüedad comunicativa, esto no ocurrió durante las entrevistas, y creemos que este tipo de errores de número no revisten gravedad.

3.1.4. Extranjerismos

Registramos un total de 875 extranjerismos no ambiguos (M = 21,87, D = 46,99), que afectaron a un total del 1,66% de las uu. ll.. Esto quiere decir que apareció un caso cada 60 unidades aproximadamente. Junto a estos que hay que considerar los 89 casos ambiguos (M = 2,22, D = 7,89), que afectaron al 0,17% de las uu. ll. (aparece uno cada 592 unidades). El valor alto de la desviación indica que los estudiantes presentan una considerable variación respecto a este error. Ciertos grupos de alumnos hicieron uso de abundantes extranjerismos, mientras otros grupos los usaron de forma puntual. La distribución por niveles revela un fuerte descenso de errores de A2 (699 más 78 ambiguos) a B1 (176 más 11 ambiguos) (tabla 8).

Nivel	No ambiguos					Ambiguos				
	Errores	Frec. rel. %	M	D	% sobre uu. ll.	Errores	Frec. rel. %	M	D	% sobre uu. ll.
A2	699	79,66%	34,95	60,62	2,63%	78	87,64%	3,90	10,89	0,30%
B1	176	20,34%	8,80	22,29	0,67%	11	12,36%	0,55	1,82	0,04%
Total	875	100%	21,87	46,99	1,66%	89	100%	2,22	7,89	0,17%

Tabla 8. Distribución por niveles de extranjerismos no ambiguos y ambiguos

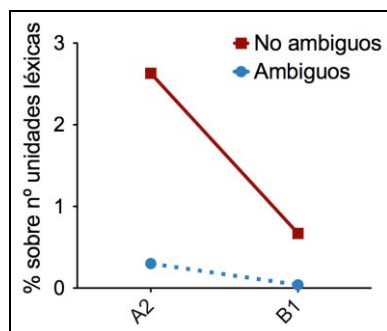


Figura 8. Extranjerismos por nivel (% sobre uu. ll. de cada nivel)

El recuento de extranjerismos no ambiguos en cada grupo se recoge en la tabla siguiente:

L1	Errores	M	D	% sobre total de uu. ll. del grupo	Un error cada...
Portugués	549	137,25	66,82	7,49%	13,35 unidades.
Alemán	116	29,00	54,02	2,54%	39,42 "
Neerlandés	71	17,75	23,07	1,37%	73,04 "
Italiano	51	12,75	5,44	0,83%	120,25 "
Inglés	30	7,50	3,32	0,63%	157,70 "
Francés	29	7,25	7,59	0,45%	221,13 "
Japonés	12	3,00	3,56	0,26%	388,66 "
Otros	9	2,25	2,63	0,19%	529,66 "
Chino	5	1,25	1,89	0,14%	718,40 "
Polaco	3	0,75	0,50	0,06%	1766,66 "

Se puede ver que lusófonos, germanófonos y neerlandeses emplearon más extranjerismos. No obstante, los altos valores de la desviación indican que tales frecuencias no son homogéneas entre los hablantes. Esto es, en dichos grupos, determinados alumnos abusaron de palabras extranjeras, mientras que otros emplearon menos. Al contrario, cometieron menos errores los hablantes de lenguas alejadas del español (japonés, chino, polaco, y las del grupo heterogéneo). Véanse ejemplos de errores no ambiguos:

- (12) AMA: [<] <las variables> de las **vogales** /// ('vocales') (PORWB1)
 (13) KRI: <es crudo / es> **liquid** /// ('liquido') (GERWA2)
 (14) MAR: <son> un poco más grande y un poco más / &eh **thicker** ... ('gruesos') (DUTWA2_2)

También recogimos 89 casos ambiguos, repartidos por grupos de L1 como muestra la tabla siguiente.

L1	Ambiguos	M	D	% sobre total de uu. ll. del grupo	Un error ambiguo cada...
Portugués	68	17,00	21,32	0,93%	107,78 unidades.
Inglés	12	3,00	3,56	0,25%	394,25 "
Chino	4	1,00	2,00	0,11%	898,00 "
Neerlandés	4	1,00	2,00	0,08%	1296,50 "
Alemán	1	0,25	0,50	0,02%	4573,00 "

En estos casos, los extranjerismos también pueden considerarse errores fonéticos, debido principalmente a la tendencia del grupo portugués a pronunciar *e* como *i* (*por iso* [por 'iso].), y en los grupos alemán, neerlandés o inglés, a pronunciar más cerrada la *e* (*es* ['is]). Véanse algunos ejemplos:

(15) *NUN: por **eso** {%pho: [por 'iso]} &mm no me {%pho: [mi]} gusta /// (PORMA2)

(16) SOR: las películas **en** {%pho: [i:n]} → el **lengua** {%pho: ['lɪŋgwa]} / original /// (ENGWB1_2)

Aunque la mayoría de préstamos proceden de la L1 de los alumnos, abundan los extranjerismos del inglés en todos los grupos, ya que los informantes son universitarios que conocen esta lengua. Algunos ejemplos son *business* o *vegetables*, o los que afectan a *en*, que resultan ambiguos porque pueden ser usos de la preposición *in* o errores de pronunciación (/en/ → [in]). Por otra parte, muchos anglicismos son términos extendidos en el lenguaje juvenil (*fast food*, 'comida rápida', *junk food*, 'comida basura', etc.).

Como indica Poulisse (1999: 55-56), el uso de una palabra de la L1 puede ser voluntario o involuntario. Por una parte, los aprendices utilizaron de manera general los extranjerismos como una estrategia de comunicación ante una dificultad de vocabulario, especialmente cuando describían una receta o los alimentos en las fotografías sobre comida: *Beijin duck*, 'pato de Pekín'; *Sauerkraut*, 'chucrut'; *sausage*, 'salchicha'; *pâte feuilletée*, 'hojaldre'; *safran*, 'azafrán'; *gamberi*, 'langostinos'; *manjeriçao*, 'albahaca', etc. Sin embargo, en otros casos el aprendiz usaba inconscientemente palabras sin traducir de su L1, seguramente por la urgencia de la comunicación oral. Así, abunda el uso de marcadores discursivos (*bon* del francés, *like* del inglés, *so* del alemán, *então* del portugués), conectores (*et* del francés, *e* del portugués y el italiano), interjecciones (*ih!* del portugués, *oh Gott!* del alemán, *oh yeah!* o *my God!* del inglés) así como fórmulas de respuesta (*ja* del alemán o el neerlandés, o *già* del italiano⁴). Más aún, algunos aprendices transfirieron directamente de su L1 palabras que guardan gran similitud formal con las correspondientes en español: verbos (*diz*, 'dice', *é*, 'es', *penso*, 'pienso'), sustantivos (*guitare*, 'guitarra'; *cirurgia*, 'cirugía'), adjetivos (*contente*, 'contento'; *difíciles*, 'difícil'), determinantes (*o*, 'el/lo'; *uma*, 'una'),

⁴ Aunque en estos casos podríamos considerar que se trata de la palabra española *ya*, creemos que es difícil que conozcan su uso como afirmación de respuesta en español. Por esta razón transcribimos estas palabras como términos extranjeros.

pronombres (*estes*, ‘estos’; *si*, ‘se’; *ce que*, ‘lo que’), numerales (*quaranta*, ‘cuarenta’; *vinte e quatro*, ‘veinticuatro’), etc. Creemos que este tipo de préstamos no suele causar dificultades en la comprensión del enunciado, a diferencia del uso consciente de extranjerismos.

Precisamente el fenómeno de transferencia de extranjerismos se agudiza entre los aprendices cuya L1 es más próxima al castellano (italianos, y sobre todo portugueses). Esto se debe probablemente a la falta de consciencia de este error, o a la creencia de que el español es una lengua fácil de aprender, por lo que existe una mayor relajación formal. Como estos préstamos son fácilmente entendidos por el oyente, no son corregidos, y su uso tiende a fosilizarse. Junto al factor de proximidad lingüística, la predisposición comunicativa del hablante puede acentuar el uso de extranjerismos, ya sea porque se atreve a hablar más (y los usa inconscientemente) o porque son una estrategia de comunicación que puede facilitar la expresión del mensaje. Esto explicaría las cifras de producción en el grupo germano o neerlandés. En ese sentido, el reducido número de préstamos en el grupo chino o japonés se puede atribuir a cualquiera de los factores anteriores. El grado de distancia tipológica entre su L1 y el español impide la transferencia directa de palabras, y análogamente, la consciencia de tener que vencer tal distancia les convierte en aprendices que tienden a controlar su expresión formal o a ser poco comunicativos. Finalmente, el grupo polaco presentó menos extranjerismos, e ignoramos las razones de estos resultados. Serían interesantes más experimentos para ahondar en la influencia de estos factores junto a otras variables.

En resumen, los extranjerismos abundaron en nuestros datos orales comparados con otros resultados obtenidos en textos escritos (Fernández López, 1990a; Fernández Jódar, 2007). La presión comunicativa de la oralidad, que carece del control formal de la escritura, parece ser una explicación plausible. Con todo, en nuestros datos, este fenómeno es muy acusado solo en determinados grupos de L1 (lusófonos, germanófonos y neerlandeses). La alta cifra de la desviación en estos alumnos indica además una elevada variación intragrupal. Por ello, serían necesarios datos de más participantes para confirmar estas tendencias. No obstante, se trata de un error fosilizable, especialmente entre nativos de lenguas afines.

3.1.5. Calcos estructurales

Los calcos se producen por la traducción palabra por palabra de una construcción de otra lengua. Aunque existen también los calcos sintácticos (p. ej., *no es posible *de entrar*, del francés *il n'est pas possible de rentrer*), en este apartado consideramos los calcos formales que afectan a unidades léxicas (nombres, locuciones, marcadores discursivos o expresiones idiomáticas). El calco del inglés *like that*, ‘así’, que los aprendices traducen por **como eso* o **como así*, nos pareció más adecuado considerarlo de nivel gramatical, y no lo incluimos entre los errores léxicos (cabe señalar que recogimos 35 casos). En esta sección se incluyen los calcos estructurales, que afectan a la forma de la palabra y dan lugar a una forma inexistente en castellano (Gómez Capuz,

2009: 7). No recogemos aquí los calcos semánticos, pues se abordan entre los errores de significado (§3.2.6), aunque los límites son borrosos en ciertos casos.

En conjunto, registramos 16 calcos estructurales ($M = 0,4$, $D = 0,84$), que afectaron al 0,03% de las uu.ll. Hubo 9 en A2 ($M = 0,45$, $D = 0,89$) y 7 en B1 ($M = 0,35$, $D = 0,81$). Véanse tres ejemplos:

- (17) GUI: he perdido mi $\rightarrow$ [/] mi hhh {%act: laugh} **carta de identidad** ... (ITAWA2) (del italiano *carta d'identità*, 'carnet de identidad')
- (18) SOR: es un mezcla de $\rightarrow$ / &eh ¿ **frutas / de mare** ? (francés *fruits de mer*, 'marisco') (ENGWB1_2)
- (19) MAR: se llama **pan de cuca** ... es como creps pero <un poco diferente> ///(DUTWA2_2) (del neerlandés *pannenkoek*, 'tortita', 'crep')

Es posible que se produjeran calcos que no identificamos como tales en el análisis, y que algunos se presten a otra interpretación (véanse los errores por perífrasis, §3.2.3, y los calcos semánticos, §3.2.6). Con todo, hay pocos casos generalizados o repetidos, y fueron causados por la prueba de descripción de fotos (p. ej., **frutas de mar*). Resulta llamativo que no se hayan producido calcos entre hablantes orientales (chinos o japoneses), lo que puede estar motivado por la mayor distancia lingüística entre su L1 y el español, o un mayor cuidado formal. En general, el aprendiz no suele ser consciente de la falta de equivalencia entre estructuras y calca directamente de su L1, aunque realmente pocos casos dificultaron la inteligibilidad del mensaje. En cambio, cuando el alumno es consciente de que el concepto no ha quedado comprendido, se suele valer de otras estrategias comunicativas (como la paráfrasis) para explicarlo.

3.1.6. Formaciones no atestiguadas en español

Recogemos en este apartado los errores léxicos por deformación, incluyendo las acuñaciones o neologismos y las distorsiones ('distortions'; James, 1998: 150). Las deformaciones por asignación incorrecta de género (p. ej., **dinera*) o de número (p. ej., [las] **Falla[s]*) no las incluimos aquí, sino en las secciones correspondientes. Ciertos casos podrían también ser clasificados en otros apartados; p. ej., **consejo*, 'aconsejo', se puede ver como confusión entre derivados de la misma raíz. En muchas ocasiones, es el nivel del alumno y su conocimiento del léxico lo que nos induce a pensar que se trata de un error formal y no una confusión entre dos palabras que conoce (p. ej., **cualidad* por 'calidad'). Asimismo, incluimos los verbos cuya raíz se desfiguró, pero siendo conjugados correctamente (p. ej., **rencomienda*, 'recomienda'). Con todo, es difícil diagnosticar ciertos casos como un problema de léxico o conjugación gramatical (p. ej., **friir* o **fritar*), o de formación de palabras (**discutiva*, 'discutible').

Según los criterios expuestos, se recogió un total de 359 errores no ambiguos de formación no atestiguada en español ($M = 8,97$, $D = 7,64$). Esto representa un 17,88% sobre el total de errores léxicos y afecta al 0,68% de las unidades léxicas (se registra un caso aproximadamente cada 147 unidades). Junto a estos hay que considerar los 133 errores ambiguos que pueden ser incorrecciones léxicas pero también fonéticas. La

distribución de estos errores muestra que disminuyen a medida que se progresa de nivel. Los alumnos de A2 cometieron 211 errores no ambiguos (el 58,77% de los 359, $M = 10,55$, $D = 8,79$) y 87 ambiguos ($M = 4,35$, $D = 5,26$). En cambio, los de B1 cometieron 148 errores no ambiguos (el 41,23% de los 359, $M = 7,40$, $D = 6,12$) y 46 ambiguos ($M = 2,30$, $D = 2,98$) (tabla 9 y figura 9).

Nivel	No ambiguos					Ambiguos				
	Errores	Frec. rel. %	M	D	% sobre uu. ll.	Errores	Frec. rel. %	M	D	% sobre uu. ll.
A2	211	58,77%	10,55	8,79	0,80%	87	65,41%	4,35	5,26	0,33%
B1	148	41,23%	7,40	6,12	0,56%	46	34,59%	2,30	2,98	0,17%
Total	359	100%	8,97	7,64	0,68%	133	100%	3,33	4,35	0,25%

Tabla 9. Distribución en niveles de errores no ambiguos y ambiguos por formación no atestiguada

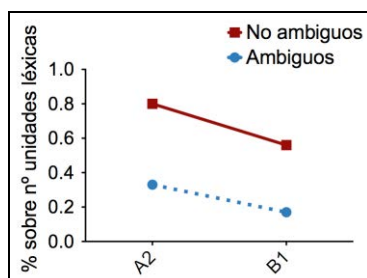


Figura 9. Errores de formación no atestiguada por nivel (% sobre uu. ll. de cada nivel)

El recuento de errores no ambiguos por grupo de L1 es el siguiente (ordenamos según tasa de error):

L1	Errores	M	D	% sobre total de uu. ll. del grupo	Un error cada...
Chino	57	14,25	3,69	1,59%	63,02 unidades.
Portugués	94	23,50	14,39	1,28%	77,97 "
Neerlandés	37	9,25	7,41	0,71%	140,16 "
Otros	34	8,50	5,26	0,71%	140,21 "
Inglés	29	7,25	4,50	0,61%	163,14 "
Japonés	24	6,00	3,37	0,51%	194,33 "
Polaco	24	14,25	2,00	0,45%	220,83 "
Francés	26	6,50	1,00	0,41%	246,65 "
Alemán	16	4,00	2,16	0,35%	285,81 "
Italiano	18	4,50	1,00	0,29%	340,72 "

Abajo se exponen ejemplos de errores no ambiguos:

- (20) ADZ: los **maloentendimentos** ... ('malentendidos') (POLMA2_1)
 (21) FIN: pues en Inglaterra estoy en **trezero** ('tercero') (ENGMB1)
 (22) LUQ: estudio español / en la facultad de / **filolofía** ('filología' o 'filosofía') (CHIMB1)

Estos errores pueden tener causa intralingüística, cuando el aprendiz realiza una falsa hipótesis sobre la forma de una palabra que no conoce enteramente. Entre los mecanismos, observamos los siguientes:

- Posible mezcla de lexemas: ej. **maraño*, 'marrón' o 'castaño'; **filolofía*, 'filología' o 'filosofía'.
- Aplicación de las reglas de derivación del castellano en casos de excepción: p. ej., **maloentendimientos* ('malentendidos'), **discutiva*, 'discutible'.
- Pérdida de consonantes (ej. **pofesora*) o vocales (**simpre*). También se produce el acortamiento de palabras por la supresión de sílaba (haplogía) en posición media (**championnes*, 'champiñones'; **pronunción*, 'pronunciación') o inicial (aféresis: **conseja*, 'aconseja').
- Suplantación de consonantes (**pepiyo*, 'pepino') o vocales (**evente*, **Guya*). En ciertos casos hay problemas fonológicos: p. ej., **incartado* ('engordado') parece causado por la no distinción entre sordas y sonoras entre los alumnos chinos.
- Adición (epéntesis) de consonantes (**escribir*, **surdeste*) o vocales (**ventidoes*, **cebollía*).
- Metátesis o cambio de orden de consonante o vocal (**trezero*, **intréprete*).
- Hipercorrección (p. ej., **aficción*), especialmente por diptongación (**entiero*). Esto es frecuente en el grupo lusófono, ya que los alumnos son conscientes del parecido formal de su L1 con el español, e hipercorrigen ciertos lexemas (Torijano, 2008). En nuestros datos hubo 13 errores solo entre los lusófonos (ej. **cierdo*, o **hacier*); en cambio, recogimos 6 en el resto de alumnos.
- A la inversa, también registramos la simplificación de diptongos (ej. **higénico*, **suficiente*).
- En otros casos, intervienen varios mecanismos de los explicados: ej. **feminimos* ('femeninos', cambio vocálico y metátesis), **tirina* ('harina', se acuña una palabra con semejanza fónica).

En otros errores por deformación rastreamos la interferencia de la lengua materna (o de otra lengua que conoce el alumno). Por una parte, pueden producirse fenómenos de relexificación (se usa la raíz léxica de una palabra de la L1 con la inflexión morfológica del español): p. ej., **crocante*, 'crujiente', del francés *croquant*. En otros casos, la forma léxica de la L1 influye de algún modo en la inflexión: p. ej., **fácilamente* (del francés *facilement*). Consideramos que al menos 86 errores de formación no atestiguada se debieron a interferencia (el 23,95% sobre un total de 359). Algunos ejemplos son los siguientes:

- (23) JUL: pide también un vino / para beber // y que está contente / con su → **escoja**
 /// (PORWA2_1) ('elección', del portugués *escolha*)

- (24) FAN: ponen los huevos / como una **omeleta** // como una → + (FREWB1)
 ('tortilla', del francés *omelette*)
- (25) AGA: [<] <patatas> / **carotas** &eh guisantes ... (ITAMA2) ('zanahorias', del italiano *carote*)

En ciertos casos es difícil rastrear la lengua que influye en la acuñación (como ocurre en **frigadora* o en **maché*), pero se intuye un mecanismo de interferencia lingüística. Cabe destacar el hecho de que muchos casos producidos por hablantes con distintas L1 se deben al inglés, como consecuencia de que es el idioma hablado por todos los participantes. Por otra parte, es significativo que este tipo de errores de interferencia apenas se haya registrado en hablantes de las lenguas más alejadas tipológicamente del español. No aparecieron en el grupo chino, y en el japonés, solamente recogimos tres casos, pero por influencia del inglés (**espisi*) y el francés (**familla*, **fácilmente*), que eran lenguas conocidas por nuestros informantes. Creemos que lo que a veces lleva a estos aprendices a una actitud poco creativa en el plano léxico puede ser la consciencia de la distancia lingüística y de la dificultad de que el interlocutor comprenda el mensaje si utiliza una raíz de su L1. Además, su tendencia a valorar negativamente el error puede producir un gran control sobre la forma lingüística. Por el contrario, el grupo lusófono registra el mayor número de formaciones no atestiguadas por interferencia de la L1, al igual que sucede con los extranjerismos. Si bien no son errores que dificulten la comprensión del mensaje, estos alumnos han de prestar una mayor atención para lograr la corrección formal sin que interfiera su lengua materna.

Por último, otros errores pueden deberse tanto a la interferencia como a mecanismos intralingüísticos. Un ejemplo es **friar*, 'freír', quizá del inglés *to fry*; **friír*, por probable influencia del italiano *friggere*; o **fritar*, por posible efecto del alemán *fritieren* (aunque también se pueden confundir formas del paradigma: *frito* > **fritar*, (yo) *frío* > **friír*). Otro ejemplo es **lemón*, del inglés *lemon*, o como deformación de *limón*. Preferimos marcar estos casos como errores de etiología desconocida.

Junto a los casos anteriores, se registraron 133 errores ambiguos de nivel fonético-fonológico o léxico. El recuento por lengua materna es el siguiente:

L1	Errores	M	D	% sobre total de uu. ll. del grupo	Un error cada...
Chino	26	6,50	4,65	0,72%	138,15 unidades.
Inglés	31	7,75	6,02	0,66%	152,61 "
Neerlandés	28	7,00	4,32	0,54%	185,21 "
Portugués	22	5,50	7,19	0,30%	333,13 "
Alemán	7	1,75	0,96	0,15%	653,28 "
Otros	7	1,75	1,71	0,15%	681,00 "
Japonés	6	1,50	1,29	0,13%	777,33 "
Italiano	4	1,00	2,00	0,07%	1533,25 "
Polaco	1	0,25	0,50	0,02%	5300,00 "
Francés	1	0,25	0,50	0,02%	6413,00 "

A continuación se transcriben algunos errores ambiguos:

(26) NUN: se llama **bacalao** {%pho: [baka'ʎao]} (PORMA2)

(27) ALE: [<] <un poquito> de / **queso** {%pho: ['keθo]} / rallado /// (ITAMB1)

La alta frecuencia de algunos de estos errores revela un fenómeno idiosincrásico restringido a aprendices aislados (ej. *bacalao* [baka'ʎao] en un alumno portugués, por interferencia fónica o de la forma *bacalhau*). Resultan más reveladoras otras incorrecciones menos numerosas pero registradas en grupos diferentes: la confusión de [s] y [θ] (*queso* ['keθo], quizá por hipercorrección), la simplificación del diptongo [ei] (*aceite* [a'sete]), la confusión de e/i (*señor* [sɪ'noɾ]), de o/u (*gusta* ['gosta]), y en menor medida, entre a/e (*sartén* [sa'tan] o [sə'ten]). En los grupos de alumnos neerlandeses, germanófonos o anglosajones, la pérdida de consonante -r al final de sílaba (*martes* ['ma:tes]) probablemente es debida a factores fonéticos; lo mismo sucede en las equivocaciones entre e/i. En muy pocos casos podemos reconocer la influencia de la L1: quizá en *oeste* ['weste] o en *suroeste* ['sur.o'weste] interfiere el inglés (*west*). La ausencia de datos escritos que confirmen si el alumno conoce la forma léxica correcta no nos permite ahondar más en la causa de estos errores, para lo cual se necesitarían pruebas complementarias.

Del conjunto de errores por formación no atestiguada, creemos que los que más dificultaron la comunicación son los que pueden llevar a la confusión entre formas muy próximas, siendo necesario el contexto para desambiguar el significado: p. ej., *bebió* [pi'bjo], puede entenderse como 'pidió' o 'vivió'; y *bebió* [bi'bio], como 'vivió'. Junto a estos, las acuñaciones a partir de una raíz de la L1 obstaculizan razonablemente la transmisión del significado (especialmente casos como **carotas*, 'zanahorias'). Por último, otra serie de neologismos resultan imposibles de relacionar con otro significante, de modo que el enunciado resulta oscuro (ej. **cado*, ¿'ubicado'?). Ante cualquiera de estas incorrecciones, los aprendices suelen servirse de estrategias de comunicación (especialmente el uso de extranjerismos, la perífrasis léxica o gestos) o de la cooperación del interlocutor para vencer sus limitaciones de vocabulario en la expresión oral. Por ello, aunque los errores formales ciertamente revisten gravedad, en la práctica es su actitud ante la situación comunicativa lo que favorece (o agrava) la comprensión del mensaje.

3.1.7. Otros errores formales

Finalmente, hubo un grupo heterogéneo de 23 errores formales que no se adscribe a ninguno de los tipos anteriores (M = 0,58, D = 1,78). Hubo 18 errores en A2 (M = 0,9, D = 2,47), y 5 en B1 (M = 0,25, D = 0,44). En conjunto afectaron al 0,04% de las uu. ll. (un caso cada 2291 unidades). Véanse ejemplos:

(28) AGA: [<] <novelas // sí> // **de** → **históricos** /// ('históricas' o 'de historia') (ITAMA2)

(29) YTO: &eh / esto / **el** / **primer lugar** / dos / chicos / &eh vienen / a → restaurante (JAPWB1_3)

- (30) JUL: creo que la gente no me / gustaba **o algo** // no sé /// (¿‘o algo así?’)
(ITAMB1)

Tres de estos de errores se explican por la mezcla de construcciones: *tiene *ganas de hambre* (‘tiene hambre’ o ‘ganas de comer’); **en generalmente* (‘en general’ o ‘generalmente’); y *novelas de *históricos* (‘novelas históricas’ o ‘novelas de historia’). Algunos podrían ser lapsus debidos a la urgencia de la comunicación oral. Otro error cometido por un aprendiz japonés se debe a la sustitución de la preposición *en* por el artículo *el* en el marcador discursivo *en primer lugar*, quizá por analogía formal con otras construcciones semejantes, o por la ausencia de preposiciones en japonés. El resto de errores afectaron a fórmulas coloquiales que expresan vaguedad o indeterminación, muy frecuentes en la oralidad, y fueron producidos por estudiantes de diferente L1: **o algo* (‘o algo así’, ‘o así’), **y todo* (‘y todo eso’, ‘y cosas así’). Creemos que estos casos revelan una competencia conversacional en desarrollo, pero carente de las formas que aportan naturalidad al discurso, quizá porque no se suele abordar su enseñanza en el aula.

3.2. Errores léxicos semánticos

Incluimos en este grupo confusiones semánticas cuya clasificación no es fácil de asignar en categorías unívocas, porque los límites de ciertos casos son borrosos. Por consiguiente, las cifras que presentamos (tabla 10 y gráfico 10) han de tomarse con cautela. Según nuestro criterio, los errores semánticos más numerosos fueron las confusiones por relación semántica: 196 errores no ambiguos (M = 4,9, D = 2,48, un 9,76% sobre el total de errores léxicos). Les siguen en número de errores las perífrasis léxicas: 54 no ambiguos (M = 1,35, D = 1,41, 2,69% sobre total de errores léxicos). Y después, las confusiones entre derivados (49 no ambiguos, M = 1,22, D = 1,35, 2,44% de errores léxicos).

ERRORES LÉXICOS SEMÁNTICOS (NO AMBIGUOS)				
	Errores	M	D	% de errores léxicos
Relación semántica	196	4,90	2,48	9,76%
Perífrasis léxicas	54	1,35	1,41	2,69%
Confusión entre derivados	49	1,22	1,35	2,44%
Otros	30	0,75	1,15	1,49%
Calcos semánticos	19	0,48	1,11	0,95%
Colocaciones	18	0,45	0,60	0,90%
Falsos amigos	17	0,43	1,01	0,85%
Registro	1	0,03	0,16	0,05%

Tabla 10. Distribución de errores léxicos semánticos

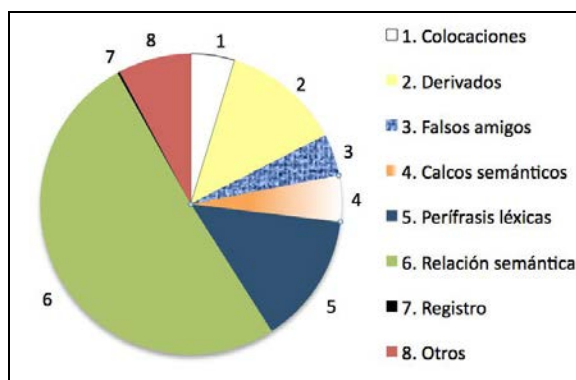


Figura 10. Distribución de errores léxicos semánticos no ambiguos (% sobre total de errores semánticos)

3.2.1. Relación semántica

Este apartado, que es el que registra mayor número de errores léxico-semánticos, incluye un conjunto de incorrecciones de distinta clase, pero en todas existe un tipo de relación semántica entre las palabras confundidas. Por una parte, abordamos aquí errores entre términos del mismo campo semántico (*salud* por *sana*), o entre los cuales existe una relación de sentido como hiperonimia o metonimia (*ensalada* por *lechuga*).

Por otra parte, examinamos incorrecciones que se deben a la neutralización de oposiciones semánticas (*ir/venir*, *ver/mirar*). Asimismo, recogemos otros errores de causa interlingüística debidos a lo que Weinreich (1953; *apud.* James, 1998: 149) denomina *polisemia divergente* o *asimetría interlingual* entre términos: las diferencias semánticas entre dos términos del español que se expresan con un mismo vocablo en la L1 del hablante (p. ej., *pedir* y *preguntar*, *to ask* en inglés). Aunque estos casos se pueden considerar *calcos semánticos*, optamos por abordarlos en esta categoría. Finalmente, también preferimos incluir aquí los errores en los que la confusión de género produce un cambio semántico (p. ej., *pimiento* frente a *pimienta*, o *bolso* frente a *bolsa*).

Conforme a estos criterios, contamos un total de 196 errores no ambiguos ($M = 4,90$, $D = 2,48$) y 8 ambiguos ($M = 0,20$, $SD = 0,52$). Respectivamente, afectaron al 0,37% y al 0,015% de las unidades léxicas. La distribución por niveles no fue demasiado desigual: en A2 se registraron 87 no ambiguos y 3 ambiguos; y en B1, 109 no ambiguos y 5 ambiguos (tabla 11 y figura 11). Parece que estas dificultades semánticas permanecen al progresar de nivel, predominando más en el nivel intermedio en nuestros datos.

Nivel	No ambiguos					Ambiguos				
	Errores	Frec. rel. %	M	D	% sobre uu. ll.	Errores	Frec. rel. %	M	D	% sobre uu. ll.
A2	87	44,39 %	4,35	2,52	0,33%	3	37,50 %	0,15	0,49	0,011%
B1	109	55,61 %	5,45	2,37	0,41%	5	62,50 %	0,25	0,55	0,019%
Total	196	100%	4,90	2,48	0,37%	8	100%	0,20	0,52	0,015%

Tabla 11. Distribución por niveles de errores de confusión entre palabras con semas comunes

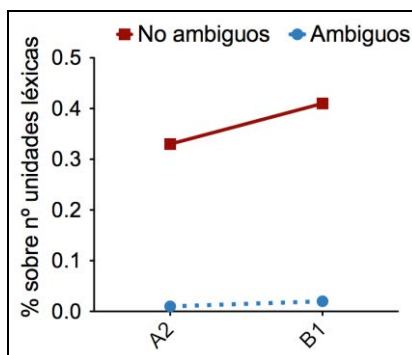


Figura 11. Confusiones entre palabras con semas comunes por nivel (% sobre uu. ll. de cada nivel)

La tabla siguiente muestra el recuento por grupo de L1 (ordenado según la tasa sobre uu. ll.).

L1	Errores	M	D	% sobre total de uu. ll. del grupo	Un error cada...
Polaco	27	6,75	2,75	0,51%	196,30 unidades.
Chino	18	4,50	2,38	0,50%	199,56 "
Neerlandés	24	6,00	2,16	0,46%	216,08 "
Alemán	21	5,25	3,20	0,46%	217,76 "
Japonés	20	5,00	2,31	0,43%	233,20 "
Inglés	17	4,25	2,63	0,36%	278,29 "
Otros	17	4,25	3,30	0,36%	280,41 "
Francés	19	4,75	2,63	0,30%	337,53 "
Italiano	16	4,00	1,15	0,26%	383,31 "
Portugués	17	4,25	3,30	0,23%	431,12 "

Algunos ejemplos de error se muestran abajo:

(31) AIS: y / como **pescas** o carne o / pollo ... ('pescado') (ENGWA2)

(32) EVE: **pidieron** a un [/] al camarero que [/] que pasó // que pasó // &ah si hay sitio para / dos personas /// ('preguntaron') (DUTWB1)

- (33) FCH: pregunta / al [/] el camarero / si hay [/] hay / &mmm dos **plazas** para / <ellos>... (DUTWA2_2) (¿‘sitio para dos?’)

Los errores entre términos del mismo campo semántico se relacionan mayoritariamente con la falta de exactitud en la designación de los alimentos (en la descripción de las fotografías): *filete* o *chuleta* para ‘carne picada’ (este caso, cometido por un estudiante polaco, es posible interferencia del término polaco *kotlety*), *chile* (‘pimiento’), *camarones* (‘langostinos, gambas’), *pimiento rojo* (‘pimentón’), *pimienta* (‘pimiento’), *pepinillos* (‘pimientos’), *almejas* (‘mejillones’), etc. En determinados errores existe incluso una relación de hiperonimia o de metonimia: *ensalada* por *lechuga* (una confusión registrada en 3 ocasiones), *carne* por ‘pollo’, *bocadillo* por ‘pan de molde’, o *conchas* por ‘mejillones’ (un error que puede estar motivado por interferencia del alemán *Muschel*). Respecto a otros campos semánticos, se registraron 5 confusiones entre términos como *pintura* o *diseño* (‘imagen’), *dibujo* (‘foto’), *cuadrado* (‘viñeta’, 2 casos) o *foto* (‘ilustración’). También hubo 4 confusiones entre términos como *bar*, *cervecería* o *tienda* para designar ‘restaurante’ (quizá por simplificación a causa de escasa concentración o nerviosismo). Algunos quizá se debieron a la estrategia comunicativa de aproximación.

Los errores por *polisemia divergente* registrados con más frecuencia fueron las confusiones entre *pedir/preguntar*: 5 casos en el grupo neerlandés; 5 en el grupo francés; 3 en el grupo italiano; 5 en el grupo polaco; 1 en un estudiante chino y 1 en una coreana. Casi todas estas incorrecciones se documentan en grupos de estudiantes cuya L1 no usa formas diferentes para ambos significados (respecto a los estudiantes coreano y chino, creemos que es interferencia del inglés). Se registraron tanto en el nivel A2 como en el B1, lo que revelaría dificultades en su aprendizaje. Otras confusiones se produjeron entre *probar/intentar* (2 en el grupo neerlandés y 1 en el polaco) y entre *tratar/probar* (1 en el grupo inglés), y las creemos debidas a la interferencia del inglés *to try*. Igualmente, el error de confusión entre *siguiente/próximo* (la **próxima historia*), registrado en una estudiante holandesa, creemos que es otro caso de interferencia, ya que el adjetivo neerlandés *volgende* equivale tanto a ‘próximo’ como a ‘siguiente’.

Respecto a la neutralización de oposiciones semánticas en los verbos, destaca la confusión entre *ir/venir*: hubo 2 casos en el grupo francés; 2 en el grupo inglés; 2 en el grupo japonés; 1 en el grupo neerlandés; 1 en una aprendiz coreana y 1 en una húngara. También fue recurrente la confusión entre *mirar/ver* (2 en el grupo italiano; 1 en el grupo francés; 1 en el grupo inglés; 1 en el grupo portugués). El resto de errores se registró en menos ocasiones: 2 casos de *llevar/traer* (en el grupo polaco) y de *saber/conocer* (en el grupo alemán), y un solo caso de otras confusiones entre *aprender/estudiar*, *andar/ir*, *buscar/encontrar*, *meter/poner*, *recordar/acordarse*, *saber/recordar* o *tardar/durar*.

Algunos de estos errores de neutralización pueden estar motivados interlingüísticamente; por ejemplo, el uso del verbo *venir* en lugar de *ir* por una alumna japonesa:

YTO: esto / el / primer lugar / dos / chicos / &eh **vienen** / a → restaurante
(JAPWB1_3)

El estudiante emplearía en japonés el verbo *kimasu*, 来ます (*kuru*, 来る), equivalente a *venir*, tomando como referencia al personaje protagonista del dibujo en lugar de distanciarse del relato, según advierte al respecto Sánchez Barrera (2009: 59). También parece interlingüística la confusión de *hablar* y *decir* propia de los aprendices brasileños (registramos 7 errores entre ellos, mientras que no aparecen en los nativos de portugués peninsular, y solamente dos en una estudiante alemana). En efecto, según indica Humblé (2005: 205), en el portugués brasileño *falar* y *dizer* tienden a mezclarse en el registro informal.

Por lo que se refiere a la neutralización del valor semántico de ciertos sustantivos, los errores más frecuentes se cometieron al usar términos como *anuncio*, *cartel*, *marca*, *señal* o *signo* para referirse a ‘letrero’ (10 casos), e igualmente, al emplear vocablos como *plaza*, *plazo* o *espacio* para aludir a ‘sitio’, lo cual se documentó en 6 ocasiones (p. ej., **plazo libre*). Especialmente, sobresale la elección de *tiempo* con el sentido de ‘horario’, ‘momento’, ‘la hora’ o incluso ‘vez’, seguramente motivado por el inglés *time* (5 ocurrencias, con la que se puede relacionar el uso erróneo de *horarios* para ‘horas’). Otros casos que nos parecen significativos son la confusión entre *pescado* y *pesca* (4 casos) y el empleo de *comidas* para aludir a ‘alimentos’ o ‘platos’ (3 ocasiones; también se registró el uso de *cocina* con el sentido de ‘comida, recetas’). Finalmente, creemos destacables, entre otros, los errores puntuales de *lengua* por *lenguaje*, *tabaco* por *estanco*, *tradición* o *hábito* por *costumbre*, *rincón* por *esquina*, *recibo* por *cuenta*, o *vestido* por *ropa* (quizá por interferencia del francés *vêtement*).

Otros errores recogidos en esta sección no muestran una relación semántica tan evidente: *estoy* **concentrándome* (‘especializándome’), **dependientes* (¿‘aduaneros?’), *son* **circular* (¿‘de forma redondeada?’), *luego la [...]* **tumbo* [*la masa de la pizza*] (¿‘la extendiendo?’), etc. También hemos de señalar el uso de *aprovechar* por *disfrutar* entre hablantes de diferente origen (un japonés, un italiano y un turco). Por el contrario, en otros casos el error viene determinado por las restricciones contextuales o sintácticas (*negó*, ‘dijo que no’; *le agradece*, ‘da las gracias’). Además, ciertos errores puntuales pueden estar inducidos por el uso erróneo del diccionario (como creemos que ocurre en *determinó*, ‘decidió’, cometido por un aprendiz japonés). Otros errores se explicarían por simplificación o uso del término más general o frecuente: **entrar a escuela*, ‘inscribirse’; *estómago* **grande*, ‘lleno’; *me vi* **buena*, ‘sana’; **comemos*, ‘cenamos’; *palabras* **malas*, ‘erróneas’, etc.

A estas incorrecciones hay que sumar los 8 errores ambiguos registrados (véase un ejemplo):

(34) MAN: me gusta → ver **película** /// (JAPWB1_3)

Dicho error es el más destacable de todos y fue producido por tres de los cuatro estudiantes japoneses, quienes emplearon la palabra *película* en singular (en lugar de plural para referirse a un conjunto de elementos indeterminados), produciendo un

error gramatical. Dicho error también puede ser léxico si se quieren referir a *cine*, una palabra con la que se confunden porque en los diccionarios ambas se traducen como *eiga* (en japonés equivalen al mismo término; Esparza Celorrio, 2007: 198). El resto de errores ambiguos son producciones idiosincrásicas que no parecen aportar mucho al análisis.

Una de las conclusiones que se puede extraer de esta relación de errores semánticos es que, pese a que existen errores recurrentes y restringidos a ciertos tipos (p. ej., *ir/venir*), la predicción de la mayoría no es sencilla debido a la capacidad creativa del estudiante, que da lugar a producciones idiosincrásicas difíciles de corregir. Aunque la comprensión no se ve afectada en ciertos contextos, en otros el significado queda envuelto en cierta vaguedad para su comprensión (p. ej., *estaba interesando*, ¿'fijándose, pendiente?'). En nuestra opinión, los errores más graves (porque dificultan más la comunicación) son los que aparecen cuando se designan realidades o entidades del mundo (p. ej., alimentos). Sin embargo, son errores que pueden registrarse en cualquier nivel, y los alumnos recurren a otras estrategias compensatorias (p. ej., la perífrasis léxica) para paliar un déficit de vocabulario. Aparte de estas carencias léxicas, los problemas relacionados con restricciones y matices semánticos pueden radicar tanto en las dificultades particulares del español como en la configuración semántica de la L1. De hecho, los casos debidos a la polisemia divergente (p. ej., la confusión *pedir/preguntar* por el inglés *to ask*) parecen resultar especialmente arduos de superar. Abordar las dificultades específicas según la L1 de cada alumno indudablemente ayudará a adquirir el componente semántico. Con todo, parece ser preciso un tratamiento didáctico más continuado a lo largo del proceso de adquisición para lograr la comprensión de ciertos matices de significado.

3.2.2. Confusión entre palabras derivadas de la misma raíz

En este apartado agrupamos los errores debidos al uso de una palabra incorrecta en lugar de otra con la misma raíz léxica. En ocasiones el diagnóstico de estos errores no es sencillo. Por ejemplo, no se clasifican en este apartado los errores de uso del adjetivo como adverbio (**normal cocino*, 'normalmente cocino') al considerarlos como errores gramaticales de estructura oracional (por cambio de función) y no errores léxico-semánticos. De igual manera, se registraron dos casos que se pueden considerar también como asignación incorrecta de género (p. ej., el verbo **cocino*, por *cocina*, 'lugar donde se cocina'), y que preferimos juzgarlos ambiguos dentro del nivel léxico.

Registramos un total de 49 errores no ambiguos por confusión entre derivados ($M = 1,22$, $D = 1,35$), junto a 2 casos ambiguos. En conjunto, afectaron al 0,09% de las unidades léxicas. La distribución por niveles es muy semejante: los aprendices de A2 cometieron 29 errores no ambiguos y 2 ambiguos; y los de B1, 20 no ambiguos (tabla 12 y figura 12). No se observa una disminución importante de casos.

Nivel	No ambiguos				Ambiguos			
	Errores	M	D	% sobre uu. ll.	Errores	M	D	% sobre uu. ll.
A2	29	1,45	1,50	0,11%	2	0,10	0,31	0,008%
B1	20	1,00	1,17	0,08%	0	0,00	0,00	0,000%
TOTAL	49	1,22	1,35	0,09%	2	0,05	0,22	0,004%

Tabla 12. Errores no ambiguos y ambiguos por confusión entre derivados

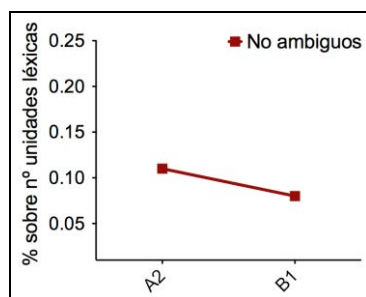


Figura 12. Errores no ambiguos de confusión entre derivados por nivel (% sobre uu. ll. de cada nivel)

Algunos errores recogidos son los siguientes:

- (35) JIY: yo como / &ah según / horario / de / **coreano** ('horario coreano' o 'de Corea') (KORWB1)
- (36) FAN: ¬ [<] <&ah> / quizás vino tinto o quizás **vino casero** hhh {%act: laugh} /// (CHIWA2_2) (¿'vino de la casa' o 'vino con casera'?)
- (37) MAK: aprendo / español / **latina** / y → griego antiguo /// ('latín') (POLWB1)
- (38) JUL: comemos menos cantidad / y → yo creo que más / &ah **variable** /// ('variado') (PORWA2_1)

Las confusiones de este tipo aparecieron entre palabras de la misma categoría (**final de semana* por *fin de semana*), a menudo cuando presentan diferente sufijo derivativo que cambia el significado (**variable* por *variado*). También fue habitual suplantar un lexema por otro de otra categoría con la que se comparte raíz: **íntimo* por *intimidación*; **latina* por *latín*; **oler* por *olor* o **alquiler* por *alquilar* (estos dos casos también podrían juzgarse como deformaciones). A este respecto, cabe reseñar las confusiones entre topónimos y gentilicios (*los platos de *chino* por 'de China'; **arroz Japón*, por 'japonés'). Dichos errores se registraron entre los aprendices asiáticos (chinos, coreanos y japoneses; sobre estos últimos, Saito, 2002: 51), aunque también se encontraron dos casos producidos por una estudiante alemana (*típico de *francés*, por 'de Francia'; *la comida de *alemán*, 'de Alemania'). En estos errores ignoramos si subyace alguna interferencia de la L1, que en otros casos sí se puede rastrear (**deportivo* por *deportista*, del italiano *sportivo*, que reúne estos dos significados del español). Con todo, la mayoría de los casos tiene causa intralingüística. Aunque ciertas incorrecciones no conllevan dificultad en la interpretación del significado (p. ej., *horario de *coreano*), otros

contextos se pueden prestar a más equívocos: *vino* *casero, por ‘de la casa’ o quizá ‘con casera’; o **al fin*, ‘por fin’ en lugar de *al final*, ‘finalmente’. Estos son precisamente los que sería importante destacar en la didáctica a cualquier grupo de alumnos.

3.2.3. Perífrasis léxica

Recogemos en esta parte errores relacionados con la estrategia comunicativa del circunloquio, que es muy productiva ante una carencia de vocabulario, y a la que los propios nativos recurrimos ante una dificultad léxica. Aunque ciertos casos se prestan a otras interpretaciones, como explicamos más abajo, juzgamos que 54 casos no ambiguos y uno ambiguo se pueden explicar por perífrasis léxica. Su distribución en niveles es muy semejante (tabla 13 y figura 13). En A2, aparecieron 25 errores (M = 1,25, D = 1,29, aproximadamente un 0,09% de uu. ll.). En B1 hubo 29 errores (M = 1,45, D = 1,54, un 0,11% de uu. ll.) y uno ambiguo (0,004% de las uu.ll.). Al igual que otros errores semánticos (como la confusión entre palabras con semas comunes o entre derivados de la misma raíz), el número de errores no disminuye de un nivel a otro de manera tan notable como ocurre con los errores de forma.

Nivel	No ambiguos			
	Errores	M	D	% sobre uu. ll.
A2	25	1,25	1,29	0,09%
B1	29	1,45	1,54	0,11%
Total	54	1,35	1,41	0,10%

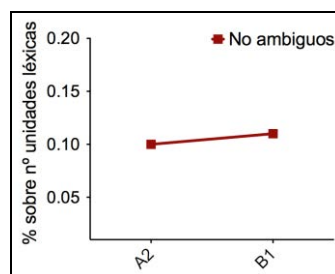


Tabla 13. Frecuencia absoluta de errores no ambiguos y ambiguos de perífrasis léxica

Figura 13. Errores de perífrasis léxica por nivel (% sobre uu. ll. de cada nivel)

Los siguientes son algunos errores etiquetados como perífrasis léxica:

- (39) AIS: lo que me encanta / es / &eh / **subir** / **montañas** (‘alpinismo’) (ENGWA2)
- (40) AMA: [<] <la> carne normalmiente / es → / **prehecha** /// (¿‘se prepara antes?’) (PORWB1)
- (41) [El alumno describe los ingredientes de la hamburguesa]
 ENT: [<] <sí el pan> suele tener /
 YTO: polvo <de hhh { %act: laugh} [/] un **polvo negro**> ///
 ENT: ¬ [<] <un [/] no / es como> [/] como una semilla // (...) que se llama sésamo /// (JAPWB1_3)

Una clase de perífrasis léxica se produce por el uso de dos lexemas para expresar el significado o el concepto que más comúnmente se expresa con uno: **árboles pera* (‘perales’), *animal* **en casa* (‘doméstico’), etc. Consideramos perífrasis de este tipo las que usan un verbo muy general como *hacer* y un sintagma nominal para expresar el

significado que normalmente se indica con un solo lexema: **hago desayuno* ('desayuno'), **hacía más nieve* ('nevaba más'), etc. En ciertos casos también se puede corregir el verbo y considerarse errores de colocación (p. ej., *me *hace confusión*, 'me confunde' o 'me produce confusión'). No obstante, preferimos ofrecer una corrección más sintética, y los incluimos en este apartado. Algunos casos también se pueden considerar calcos de otra lengua que conoce el hablante: p. ej., **fago confusión* puede ser calco del inglés *I make a mistake*; **he hecho una cirugía* ('me operé'), del inglés *I have had a surgery*; *productos *de leche* ('lácteos'), del inglés *milk products* o el alemán *Milchprodukte*; **comer la cena* ('cenar'), del inglés *eat the dinner* o el alemán *Iss das Abendessen*; y **tienda de hamburguesas* ('hamburguesería'), del inglés *burger shop*.

Otro tipo es la denominada perífrasis descriptiva, cuando el no nativo emplea palabras que describen características o rasgos del lexema reemplazado: *un *polvo negro* ('sésamo'), **caja para llevar* ('tartera'), *comida muy *fácil* (¿'rápida?'), *palabras *malas* ('tacos, palabras groseras'). Otras explican el concepto o el objeto referido: *como un *tablero* ('una bandeja'), *es como los *suelos de cerdo* ('pezuñas'), etc.

Por último, recogimos otras perífrasis que se pueden estimar como de aproximación, en los que la construcción usada expresa un sentido próximo e impreciso respecto a la construcción más apropiada: **llaman a comer para dos* ('piden una mesa para dos'); *no hay *libres lugares* ('sitio' o 'mesa libre'); *la carne es *prehecha* ('se prepara antes' o 'ya está preparada'); *hablaba [...] *con espacio* (¿'despacio' o 'con pausas?'), etc. Este último tipo de perífrasis es menos fácil de juzgar, pues en ciertos contextos causan un problema de adecuación y no tanto de agramaticalidad. Tampoco está clara la clasificación de ciertos errores, que se pueden considerar también de relación semántica.

Junto a los anteriores, recogimos otro en una alumna japonesa que consideramos ambiguo porque puede interpretarse también como un error de estructura oracional por omisión de un sustantivo:

- (42) JTO: ¿ envuelvo [/] &vo con // <transparente>? (JAPWB1_3) ('papel' o 'film transparente de cocina')

Para concluir, es preciso señalar la dificultad de clasificar este tipo de errores. En ciertos contextos podemos advertir la influencia de la lengua materna y pueden ser también calcos (p. ej., **árboles manzana*, 'manzanos', del equivalente chino 苹果树, que utiliza varios caracteres para expresar el concepto). Dado que los errores recogidos son puntuales, debido a la escasez de participantes de cada L1, no se puede confirmar si se trata de casos que revelan tendencias de error de determinados grupos de alumnos. En otros contextos, parece ser un mecanismo intralingüístico, muy idiosincrásico, que facilita la comunicación. Con todo, independientemente de que se trate de uno u otro tipo de error, en limitadas ocasiones dificulta la comprensión del mensaje.

3.2.4. Colocaciones

Los casos que clasificamos como errores de colocación son en realidad errores por relación semántica entre la palabra equivocada y la correcta, la cual solamente se utiliza en determinadas combinaciones con ciertos términos (Higueras García, 2006). Por ejemplo, *dar* y *dispensar* comparten el sema de OFRECER o ENTREGAR, pero el uso de *dispensar* se restringe a palabras como ‘honor’, ‘interés’ o ‘admiración’. Por ello, algunos casos que se reúnen en esta sección bien podrían aparecer en el apartado de relación semántica, aunque nos parece que más bien es la frecuencia combinatoria fijada por el uso (y no tanto las diferencias o matices en el significado) lo que causa los errores aquí recogidos.

Contabilizamos 18 colocaciones incorrectas en nuestro corpus (M = 0,45, D = 0,60), de las cuales 10 se registraron en A2 (M = 0,5, D = 0,5) y 8 en B1 (M = 0,4, D = 0,5) (tabla 14 y figura 14). En conjunto, solo afectaron al 0,034% de las uu. ll. (aparecería un error casi cada 2927 unidades). Teniendo en cuenta el escaso número de errores de este tipo, al observar su distribución por niveles no podemos deducir que tienden a desaparecer en niveles más altos, aunque la tasa más baja de error corresponda a B1.

Nivel	No ambiguos			
	Errores	M	D	% sobre uu. ll.
A2	10	0,50	0,69	0,038%
B1	8	0,40	0,50	0,030%
Total	18	0,45	0,60	0,034%

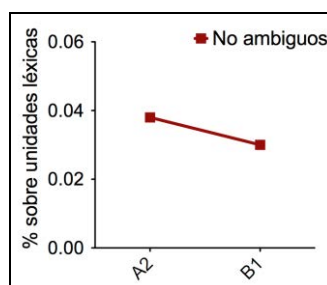


Tabla 14. Errores no ambiguos de colocaciones

Figura 14. Colocaciones erróneas por nivel

Los siguientes son ejemplos de colocaciones erróneas:

(43) THO: ¬ [<] <como> un bocadillo o un / **parte** de <pizza> (‘trozo’ o ‘porción’) (GERMA2)

(44) JUS: robar / **usar** <drogas> (‘consumir’) (PORWA2_2)

El error en una colocación se puede producir en la elección de la palabra *nodo* o *base* (un verbo o un sustantivo generalmente) o del *colocativo* que lo selecciona (respectivamente, un sustantivo o adjetivo). En ciertos casos los errores están motivados por interferencia lingüística (se pueden considerar calcos): p. ej., *vino* *rojo, del inglés *red wine* o el alemán *Rotwein*; *pago *atración (‘atención’), del inglés *pay attention*; *usar drogas, del inglés *use drugs*. En otros errores hay simplificación de estructuras por el uso de una palabra más general: *lluvia* *pequeñita, *tiempo* *grande, *tenía una carrera, *hizo un festivo. En todo caso, se trata de incorrecciones que no impiden la comprensión del significado, pero producen un enunciado inadecuado («suena raro»), y evidencian una falta de precisión respecto al contexto. Los errores en las colocaciones

son algunos de los que más contribuyen a identificar al hablante como extranjero, y deben corregirse especialmente cuanto más se progresa de nivel.

3.2.5. Falsos amigos

Aunque también pueden ser considerados errores léxicos formales (James, 1998: 147), incluimos los falsos amigos entre las incorrecciones semánticas, siguiendo a Fernández López (1990a). Realmente, son casos específicos de calcos semánticos (Gómez Capuz, 2009), pero optamos por reunirlos aquí por su frecuencia de aparición y su carácter menos idiosincrásico. En total, recogimos 17 errores no ambiguos ($M = 0,43$, $D = 1,01$), que afectaron al 0,03% de las uu. ll. Es posible que algún error de interferencia recogido en otras secciones se pueda considerar junto a estos casos. Mostramos abajo algunos ejemplos:

- (45) AMA: puedo **ordenar** un [/] &mm uno de los platos ('pedir'; de *to order*) (PORWB1)
- (46) MSU: **realicé** que → / había (...) [/] me lo había olvidado ('me di cuenta', de *to realize*) (JAPWB1_2)

La aparición de falsos amigos por niveles fue de 6 en A2 ($M = 0,30$, $D = 0,92$) y 11 en B1 ($M = 0,55$, $D = 1,10$). Esta limitada cifra de errores impide extraer conclusiones acerca de su relación con el nivel de competencia. Por otro lado, ciertos casos se registraron repetidamente en un mismo hablante: *cierta* (por 'correcta'), 3 casos en el hablante PORWA2_2; *ordenar* (por 'pedir'), 3 casos en PORWB1; *en fin* (por 'finalmente')⁵, 3 casos en ITAMB1; o *memoria* (por 'recuerdo'), 2 casos en JAPWB1_2. Esta distribución de errores restringida a ciertos estudiantes revela que los falsos amigos pueden ser un fenómeno con tendencia a la fosilización en determinados aprendices.

Algunos falsos amigos procedentes del inglés se registran en grupos de diferente lengua materna (p. ej., *ordenar*, de *to order*, por 'pedir', o *realizar*, de *to realize*, por 'darse cuenta'), dado que el inglés es la segunda lengua de casi todos los aprendices participantes. Según nuestro criterio, creemos que los falsos amigos no constituyen errores abundantes, pero su importancia es considerable porque pueden cambiar ampliamente el significado del mensaje. Por ello, cada grupo específico de estudiantes ha de ser advertido acerca de los términos que pueden ser problemáticos por interferencia de su lengua materna.

3.2.6. Calcos semánticos

Según Gómez Capuz (2009: 8), en los calcos semánticos se traduce un término o expresión extranjera mediante palabras o construcciones existentes en castellano, pero

⁵ El caso de *en fin* también puede interpretarse como una confusión entre dos locuciones: *al final*, 'finalmente', y *en fin*, 'en resumidas cuentas' (sería intralingüístico). Para decidir en este tipo de disyuntivas, ha primado la hipótesis de la influencia de la L1.

no con el significado propio o más adecuado en español, sino con el de la lengua de la que se traducen. Por ejemplo, *puedo hablar inglés* es una traducción de *I can speak English*, que en castellano equivale a ‘sé hablar’ o ‘hablo inglés’ (el error se produce respecto al significado expresado, pues la forma es correcta). En conjunto, estimamos que hubo 19 errores ($M = 0,48$, $D = 1,11$), que representaría un total del 0,04% de las uu. ll. (un error casi cada 2773 unidades). Hubo 8 casos en A2 ($M = 0,4$, $D = 0,82$) y 11 en B1 ($M = 0,55$, $D = 1,36$). Con todo, ciertos casos son difíciles de distinguir de los calcos estructurales. Véanse los siguientes ejemplos:

- (47) GIU: le gusta salir le gusta **hacer fiesta** cada día (it. *fare festa*, ‘salir de fiesta’) (ITAWA2)
- (48) REM: **la mayoría del tiempo** los españoles no hacen como muchas errores (FREMB1) (‘la mayoría de las veces’, del francés *la plupart du temps*)
- (49) FIN: he vivido / con mucho [/] muchos jóvenes / y creo que / hay muchos que no **pueden** & coc [/] cocinar (‘saben’, del inglés *they cannot cook*) (ENGMB1)

El calco registrado en más grupos de L1 es *poder* + infinitivo con el sentido de *saber*, por influencia del inglés. Aparecen otros casos, pero son puntuales: en **son menos prisas* (‘tienen menos prisa’) creemos que influye la estructura del inglés (*they are less hurried*); y el error de *la mayoría del *tiempo* se repite tres veces en un solo aprendiz (lo cual puede mostrar una tendencia hacia la fosilización). En cambio, el error de **hacer fiesta* (‘salir de fiesta’) fue cometido por distintos aprendices italianos y franceses, y en consecuencia es probable registrarlo con más hablantes de dichas lenguas maternas.

Además de los anteriores, recogimos tres casos de *encontrar* con el sentido de ‘conocer’ o ‘encontrarse con’ o, en el grupo francés (uno en A2, y dos en B1) (véase un ejemplo):

- (50) [Da un consejo a un estudiante para que mejore su nivel de español]
 FAN: y **encontrar** a gente española // porque es difícil cuando → es [/] <eres>
 ENT: [<] <claro> ///
 FAN: ¬ erasmus / **encontrar** gente española /// (FREWB1)

Se trata de un error ambiguo que se puede estimar como un problema gramatical de omisión del pronombre *se*, o un error léxico por uso de *conocer* (de *rencontrer des espagnols*, calco o falso amigo).

3.2.7. Registro

Los errores de registro afectan realmente al nivel pragmático de la lengua, pero aquí abordamos solo impropiedades léxicas relacionadas con el grado de formalidad de la comunicación. En cualquier caso, apenas marcamos errores por registro en nuestros datos. Una de las razones se debe al propio método de obtención de datos, la entrevista oral, que se desarrolló en un entorno y una situación comunicativa informal. Por ello, dudamos si considerar incorrectos usos léxicos como los siguientes.

(51) [En la prueba de descripción de la historia a partir de las viñetas]

FAN: parece feliz **el tío** <hhh {%act: laugh}> /// ('el hombre') (FREWB1)

(52) FCH: la gente → &vie [/] **viejos** ('mayor' o 'de la tercera edad') (FREMA2)

El caso de *gente vieja* puede ser despectivo en determinados contextos, pero creemos que no ocurre esto en el enunciado del ejemplo. Aunque dudosos, preferimos no estimar erróneos estos usos para evitar la hipercorrección. Así, se marcó solo un error de registro (M = 0,03, D = 0,16), que es el siguiente:

(53) JUS: las cosas que hhh {%act: laugh} [/] que → [/] que [/] que / **acontecía** conmigo ('ocurrían') (PORWA2_2)

La causa de este error, no obstante, se puede explicar por la influencia del portugués, lengua con la que el español comparte muchísimo vocabulario, pero de uso diferente en cada lengua respecto al registro o los matices semánticos específicos. Para profundizar en el estudio de la variación de registro por parte de los aprendices, sería necesaria la ampliación de nuestros datos con otras pruebas diseñadas específicamente a tal efecto, ya que el diseño de nuestra entrevista no fue concebido con ese propósito.

3.2.8. Otros errores semánticos

Reunimos finalmente en este apartado un conjunto de errores semánticos que no pudimos adscribir a ninguno de los anteriores. En algunos no encontramos relación semántica evidente entre las palabras confundidas, pero tampoco nos pareció que el alumno usara una perífrasis léxica. En otros errores tampoco pudimos conocer el mensaje que pretendía expresar el hablante porque carecemos de contexto lingüístico (en estos casos, no ofrecemos corrección). Asimismo, cuatro errores aquí incluidos se produjeron porque el estudiante repite las palabras del entrevistador sin comprender lo que expresa (véase un ejemplo):

ENT: [<] <le pide> / consejo ¿no?

FRA: ¡ah! sí /// **le@g pide@g consejo@g de** → / tomar / pollo ('le aconseja') (GERWB1_1)

En total, recogimos 30 errores no ambiguos (M = 0,75, D = 1,15). Su distribución por niveles fue muy parecida: 16 no ambiguos en A2 (M = 0,8, D = 1,15) y 14 en B1 (M = 0,7, D = 1,17). En conjunto, afectaron al 0,06% de todas las uu. ll. (un error casi cada 1756 unidades). Véanse ejemplos:

(54) ALE: **nos quedamos** muy bien todos &l [/] los del pisos /// ('nos llevamos') (ITAMB1)

(55) ADC: también **trato** Erasmus → / como una oportunidad (¿'considero?') (POLMB1)

(56) REM: y **se fue** el teléfono ///

ENT: <hhh {%act: laugh}> ///

REM: [<] <pero> funciona ahora /// (¿'se estropeó'?) (FREMB1)

Algunos errores de este apartado se produjeron por simplificación del vocabulario y uso de verbos demasiado generales: *se *fue el teléfono* ('se estropeó'), **coger la estructura de la lengua* ('adquirir', 'aprender'), *el camarero que *pasó* (¿'se acercó?'). Por último, ciertos casos no obedecen a ninguna causa explicable y podrían ser lapsus debidos a falta de concentración o a la urgencia de la oralidad: *el camarero *piensa* (*piensa*, 'dice') *que...; estudio inglés // ¿hace cuánto *estuve?* (¿'empecé?'). También recogimos un error por el uso de *por eso* en lugar de la fórmula de respuesta *eso es*, ya que el aprendiz lo empleaba como marcador polifuncional. No pudimos explicar el resto de casos porque no fue posible reconstruir la intención comunicativa del hablante (p. ej., *la cara del camarero está muy *mal*).

Por último, recogimos tres errores ambiguos en el nivel A2. Fueron producciones idiosincrásicas por repetición mecánica de palabras, lapsus y fallos de concentración, o por la simplificación de estructuras.

4. Conclusiones sobre los errores léxicos

Nuestra aportación al estudio de los errores léxicos en la oralidad es la evidencia de que estos reflejan características propias de la producción oral, y por tanto sus resultados pueden diferir de los del análisis de la escritura. Por ello, el estudio de ciertos errores de la oralidad quedaría más completo si se complementara con otras pruebas (p. ej., escritas o longitudinales), para confirmar ciertas tendencias de error, o si estos se deben a la competencia o la actuación (p. ej., debido a la falta de concentración). Este es uno de los aspectos mejorables de nuestro trabajo, junto al hecho de que los datos que analizamos son limitados, por lo que se refiere al escaso número de alumnos de cada grupo de lengua materna, así como el nivel de competencia, restringido a intermedio-bajo. Sería también positivo validar la taxonomía y la anotación de los errores de manera independiente entre varios etiquetadores, para valorar la consistencia del etiquetado y su posible extensión a otros corpus. La dificultad de clasificar determinados errores obliga a interpretar nuestros resultados con prudencia. Con todo, las principales incorrecciones que los estudiantes de español presentaron en nuestros datos se pueden resumir en los siguientes puntos:

- Destaca el progreso positivo respecto a los errores formales, como confirman nuestros datos comparando los de A2 con el menor número en B1, lo cual es un hecho ya constatado en otros análisis de errores de la escritura (Vázquez, 1999).
- Algunos errores formales no desaparecen desde los primeros niveles y pueden fosilizarse si no se abordan desde el inicio; p. ej., el uso de *gente* con complementos en plural o las interferencias léxicas de las que el alumno no es consciente, sobre todo entre aprendices cuya L1 es el portugués.
- Los errores semánticos, aun registrados con una frecuencia menor, muestran una dificultad más persistente en el progreso de adquisición. La necesidad de rigor a medida que se progresa de nivel impone al alumno ser más consciente

de restricciones y matices semánticos, pero este progreso parece verse obstaculizado por la configuración semántica de la L1 o por las dificultades del español.

- La urgencia de la comunicación o la relajación formal debida al contexto informal de la entrevista puede explicar la profusión de extranjerismos (especialmente entre lusófonos), o incluso la producción de malapropismos, que los propios nativos cometemos en momentos de distracción.
- Asimismo, los errores fonéticos pueden crear ambigüedad para comprender el mensaje ([*ˈotɾəs*], ¿*otros* u *otras*?), y el análisis exclusivo de los datos orales a veces no basta para diagnosticar el error.
- Las diferencias particulares entre alumnos de grupos de diferente L1, aunque no son generalizables, en ciertos casos pueden revelar un mayor control sobre la producción o mayor censura del error (p. ej., entre los alumnos chinos o japoneses), frente a otros que parecen sentirse más cómodos o confiados expresándose en español, pero con menor rigor formal (p. ej. los aprendices lusófonos).

Respecto a las implicaciones didácticas, en ciertos casos parece necesario abordar un tratamiento didáctico específico considerando las dificultades de cada grupo de aprendices según su L1 (p. ej., respecto a los calcos, los fenómenos semánticos de polisemia divergente o los falsos amigos, así como los errores debidos al parecido formal entre el español y otras lenguas románicas). En líneas generales, los errores formales precisan una corrección desde los primeros niveles para evitar su fosilización. Respecto a la adquisición del componente semántico de sustantivos, verbos o adjetivos, se requiere una insistencia constante y prolongada en el aprendizaje (especialmente cuando puede existir interferencia). Sin olvidar que la adquisición del léxico de cualquier lengua es una tarea inacabada y en continua reactivación.

Recibido: 24.01.2013

Aceptado: 15.05.2014

Referencias bibliográficas

- Baralo, M. (1996): «Algunos aspectos de la adquisición de la morfología léxica del español como lengua extranjera». En Montesa, S. y P. Gomis, (eds.). *Actas del V Congreso Internacional ASELE*. Málaga, Univ. de Málaga, vol. I, págs. 143-150. cvc.cervantes.es/ensenanza/biblioteca_ele/antologia_didactica/morfologia/baralo.htm (26/04/14)
- Baralo, M. (2001): «El lexicón no nativo y las reglas de la gramática». En Pastor, S. y V. Salazar (eds.): *Tendencias y líneas de investigación en adquisición de segundas lenguas*. Alicante, Universidad de Alicante, págs. 23-38: http://cvc.cervantes.es/ensenanza/biblioteca_ele/antologia_didactica/morfologia/baralo_01.htm (26/04/14).
- Barlow, M. (2005): «Computer-based analyses of learner language». En Ellis, R., y G. Barkhuizen, *Analysing Learner Language*, cap. 14, Oxford, Oxford University Press, págs. 335-357.

- Benedetti, A. M^a. (2002): «El portugués y el español frente a frente: aspectos fonético-fonológicos y morfosintácticos», *Carabela*, 51 (Monográfico: Lingüística contrastiva (I)), Madrid, S.G.E.L., págs. 147-171.
- Campillos Llanos, L. (2012): *La expresión oral en español lengua extranjera: interlengua y análisis de errores basado en corpus*. Tesis doctoral inédita. Universidad Autónoma de Madrid. Dpto. de Lingüística.
- Collentine, J. (2004): «The Effects of Learning Contexts on Morphosyntactic and Lexical Development», *Studies in Second Language Acquisition*, 26(2), págs. 227-248.
- Consejo de Europa (2001): *Common European Framework of Reference for Languages: Learning, Teaching, Assessment*. Cambridge, Cambridge University Press.
- Corder, P. (1971/1991): «Idiosyncratic dialects and error analysis». *International Review of Applied Linguistics*, 9(2), págs. 147-60. Traducción al español: «Dialectos idiosincrásicos y análisis de errores», en Muñoz Licerias, J. (comp.) (1991): *La adquisición de las lenguas extranjeras*, Madrid: Visor, págs. 63-77.
- Cresti, E., y M. Moneglia (2005): *C-Oral-Rom. Integrated Reference Corpora for Spoken Romance Languages*. Amsterdam, John Benjamins. Studies in Corpus Linguistics, 15.
- Cruz Piñol, M. (2012): *La lingüística de corpus y la enseñanza del español como lengua extranjera*. Madrid: Arco/Libros, S.L.
- Dagneaux, E., S. Denness, y S. Granger (1998): «Computer-aided error analysis», *System*, 26(2), págs. 163-174.
- Díaz-Negrillo, A., y J. Fernández Domínguez (2006): «Error tagging systems for learner corpora», *RESLA*, 19, págs. 83-102: http://dialnet.unirioja.es/servlet/dfichero_articulo?codigo=2198610&orden=72810 (26/04/14)
- Díaz Rodríguez, L. (2007) *Interlengua española: estudio de casos*. Barcelona, Printulibro Intergrup.
- Esparza Celorrio, S. (2007): «La expresión escrita en español de los estudiantes japoneses: 11 errores frecuentes que deben tratarse en clase», *Journal of Enquiry and Research*, 85, Kansai Gaidai University, págs. 195-208.
- Fernández Jódar, R. (2007): *Análisis de errores léxicos, morfosintácticos y gráficos en la lengua escrita de los aprendices polacos de español*. Tesis doctoral. Universidad Adam Michiewicz, Poznan, Polonia, www.mecd.gob.es/redele/Biblioteca-Virtual/2007/memoriaMaster/2-Trimestre/FERNANDEZ-J.html (26/04/14)
- Fernández López, S. (1990a): *Análisis de errores e interlengua en el aprendizaje del español como lengua extranjera*. Tesis doctoral. Universidad Complutense de Madrid. Fac. de Filología, Dpto. de Filología Románica.
- Fernández López S. (1990b): «Formación de palabras y adquisición de la lengua». En *Actas del VIII Congreso Nacional de AESLA (Vigo, 1990)*, Vigo, Universidad de Vigo. págs. 255-265. http://cvc.cervantes.es/ensenanza/biblioteca_ele/antologia_didactica/morfologia/sn_fernandez.htm (26/04/14)
- Fernández Silva, C. (1998): «La creación léxica en la interlengua de español». En Celis, A., y J. R. Heredia (coords.): *Actas del VII Congreso de ASELE (Almagro, 1996)*. Cuenca: Ed. Universidad de Castilla-La Mancha, págs. 213-219. http://cvc.cervantes.es/ensenanza/biblioteca_ele/asele/pdf/07/07_0211.pdf (26/04/14)
- Gómez Capuz, J. (2009): «El tratamiento del préstamo lingüístico y el calco en los libros de texto de bachillerato y en las obras divulgativas», *Tonos. Revista Electrónica de Estudios*

- Filológicos*, XVII. www.tonosdigital.es/ojs/index.php/tonos/article/viewFile/294/203 (26/04/14)
- Granger, S. (2003): «Error-tagged Learner Corpora and CALL: a promising synergy», *CALICO Journal*, 20(3), págs. 465-480.
- Granger, S. (2012): «How to use foreign and second language learner corpora». En Mackey, A., y S. M. Gass (eds.): *Research methods in Second Language Acquisition: A practical guide*. London, Blackwell Publish, págs. 7-29.
- Hernández Fernández, A. (2004): *Los errores lingüísticos*. Valencia, Nau Llibres.
- Higueras García, M. (2006): *Las colocaciones y su enseñanza en la clase de ELE*. Madrid, Arco/Libros, S.L.
- Humblé, P. (2005): «Falsos cognados. Falsos problemas. Un aspecto de la enseñanza del español en Brasil», *Revista de Lexicografía*, 12, págs. 197-207: http://ruc.udc.es/dspace/bitstream/2183/5510/1/RL_12-7.pdf (26/04/14)
- James, C. (1998): *Errors in Language Learning and Use. Exploring Error Analysis*. London/New York, Longman. Applied Linguistics and Language Study Series.
- Lafford, B. (2004): «The Effect of the Context of Learning on the Use of Communication Strategies by Learners of Spanish as a Second Language», *Studies in Second Language Acquisition*, 26(2), págs. 201-226.
- Lafford, B. y J. G. Collentine (1987): «Lexical and Grammatical Access Errors in the Speech of Intermediate/Advanced Level Students of Spanish», *Lenguas Modernas*, 14, págs. 87-112.
- Liang, W.-F., y M. Llobera (1998): «Estrategias de comunicación en el interlenguaje del español de los hablantes de chino mandarín». En Celis Sánchez, M^a. A., y J. Ramón Heredia (coords.): *Actas del VII Congreso de ASELE*, págs. 291-300. http://cvc.cervantes.es/ensenanza/biblioteca_ele/asele/pdf/07/07_0289.pdf (22/04/14)
- Lozano, C. 2009: «CEDEL2: Corpus Escrito del Español como L2». En Bretones, C. M., *et alii* (eds.): *Applied Linguistics Now: Understanding Language and Mind/La Lingüística Aplicada actual: Comprendiendo el Lenguaje y la Mente*. Almería, Universidad de Almería, págs. 197-212.
- Marsden, E. y A. David (2008): «Vocabulary use during conversation: a cross-sectional study of development from year 9 to year 13 among learners of Spanish and French», *Language Learning Journal*, 36(2), págs. 181-198.
- McCarthy, C. (2008): «Morphological variability in the comprehension of agreement: an argument for representation over computation». *Second Language Research*, 24(4), págs. 459-486.
- Mitchell, R., L. Domínguez, M^a. J. Arche, F. Myles, y E. Marsden (2008): «SPLLOC: A new database for Spanish second language acquisition research». En Roberts, L., F. Myles y A. David (eds.): *EuroSLA Yearbook*, 8, Amsterdam/Philadelphia, John Benjamins, págs. 287-304.
- Moreno Sandoval, A., y J. M. Guirao (2006): «Morpho-syntactic Tagging of the Spanish C-ORAL-ROM Corpus: Methodology, Tools and Evaluation». En Kawaguchi, Y., S. Zaima y T. Takagaki (eds.): *Spoken Language Corpus and Linguistic Informatics*, Amsterdam, John Benjamins, págs. 199-218.

- Muñoz Benito, A. M^a. (2002): *Estrategias de comunicación en la interacción oral entre hablantes nativos y no nativos de español: Implicaciones y propuestas didácticas*. Tesis doctoral. Universidad Autónoma de Madrid.
- Pinilla, R. (2000): «El desarrollo de las estrategias de comunicación en los procesos de expresión oral: un recurso para los estudiantes de ELE», *Carabela*, 47, págs. 53-67.
- Poulisse, N. (1990): *The Use of Compensatory Strategies by Dutch Learners of English*. Pordrecht, Foris Publ.
- Rubio, F. (2005): «El uso de estrategias comunicativas entre hablantes avanzados de español». En Álvarez, A., et alii (2006): *Actas del XVI Congreso Internacional de ASELE*, Oviedo, 22-25 de septiembre de 2005, págs. 547-556. cvc.cervantes.es/ensenanza/biblioteca_ele/asele/pdf/16/16_0545.pdf (26/04/14)
- Saito, A. (2002): *Análisis de errores en la expresión escrita de los estudiantes japoneses*. Memoria de máster. Universidad de Salamanca. www.educacion.es/redele/biblioteca2005/saito/saito.pdf (26/04/14)
- Sánchez Barrera, M. (2009): «Análisis de errores en la prueba de expresión oral del examen D.E.L.E inicial», *Boletín del Dpto. de Español de Univ. Kanagawa*, págs. 49-65. <http://klibredb.lib.kanagawa-u.ac.jp/dspace/bitstream/10487/5442/1/%E8%A8%80%E8%AA%9E%E7%A0%94%E7%A9%B6%E3%80%8049-65.pdf> (26/04/14)
- Tarone, E. (1980): «Communication strategies, foreigner talk and repair in interlanguage», *Language Learning*, 30, págs. 417-431.
- Torijano Pérez, J. A. (2008): «El estudio de los determinantes en aprendices lusohablantes de español», *DICENDA. Cuadernos de Filología Hispánica*, 26, págs. 235-257.
- Vázquez, G. (1999): *¿Errores? ¡Sin falta!* Madrid, Edelsa.
- Whitley, S. (2004): «Lexical Errors and the Acquisition of Derivational Morphology in Spanish», *Hispania*, 87(1), págs. 163-172. <http://www.jstor.org/stable/20063018> (22/04/14)
- World Wide Web Consortium (W3C) (2012): «Extensible Markup Language (XML)»: <http://www.w3.org/XML/>