

ARTICLE

The role of audience design and goal bias in message generation: Evidence from Chinese source-goal motion events

Chen Zhao, Rui Xu and Tingting Sun

Bilingual Cognition and Development Lab, Center for Linguistics and Applied Linguistics, Guangdong University of Foreign Studies, Guangzhou, China

Corresponding author: Chen Zhao; Email: zhao_chen001@sina.com

(Received 28 March 2024; Revised 15 May 2025; Accepted 08 July 2025)

Abstract

This study investigated the effects of audience design and goal bias in Chinese speakers' message generation of source-goal motion events (e.g., *A bird flies from the tree to the house*), using picture description and memory tasks. The status of the source (e.g., the tree) or the goal (e.g., the house) was manipulated as known or unknown to the confederate addressees. The findings revealed that the participants were more likely to omit the sources when they were mutually known to the addressee than when they were not. However, participants showed similar accuracy in detecting source changes, regardless of whether the sources were known to the addressee. Moreover, they consistently mentioned goals and showed similar accuracy in detecting goal changes, regardless of whether the goals were known or unknown to the addressee. The results suggest that audience design influenced the speakers' mention of sources, but not their memory of them. It did not affect either the mention or the memory of goals. Goal bias was not consistently observed across the two experiments, both linguistically and in memory. This suggests a fragile goal bias in Chinese. Taken together, these findings demonstrate that audience design and goal bias influence the message generation of motion events in Chinese speakers.

Keywords: audience design; goal bias; language production; message generation

1. Introduction

Language production mainly consists of three stages: message planning, linguistic formulation and articulation (Levelt, 1989). Message generation, also called message planning or conceptualization, refers to the process by which the speaker encodes the conceptual contents they prefer to express. It functions as the input to the downstream processes of language production (Bock & Levelt, 1994; Levelt, 1989).



Source-goal motion events are often taken as a testing ground to investigate message generation processes due to their well-studied conceptual structures (Papafragou & Grigoroglou, 2019). Generally, a source-goal motion event is composed of an object (i.e., the figure) moving from a starting point (i.e., the source) to an endpoint (i.e., the goal) (Talmy, 1985). For example, in the source-goal motion event '*The bird flies from the tree to the house*', '*the bird*' is the figure, '*flies*' is the manner, '*the tree*' is the source and '*the house*' is the goal. When speakers talk about a motion event, they may only mention the goal '*The bird flies to the house*', with the source (*the tree*) omitted; or only mention the source '*The bird flies away from the tree*' with the goal (*the house*) omitted or mention both event elements like '*The bird flies from the tree to the house*'. The flexibility for speakers to mention either the source, the goal or both the source and the goal in motion events enables us to examine speakers' selective choice of different event elements like the source or the goal during the message generation processes.

Many factors may affect the processes of message generation (see Konopka & Brown-Schmidt, 2014; Papafragou & Grigoroglou, 2019 for review). For example, between-language differences (Liao et al., 2020; Slobin, 1996; von Stutterheim & Nüse, 2003), cognitive factors like goal (over source) bias, i.e., the high prominence in conceptual representation of endpoints over starting points of a motion event (Do et al., 2020, 2022), audience design (Horton & Gerrig, 2002, 2016; Horton & Keysar, 1996) and discourse context such as visual context (Brown-Schmidt & Konopka, 2008; Brown-Schmidt & Tanenhaus, 2006) and linguistic context (Brennan & Clark, 1996; Van Der Wege, 2009).

Among these factors, the effect of audience design and goal bias on message generation is still controversial. Audience design refers to the cooperative process in language production that speakers formulate utterances based on the knowledge state or informational needs of their listeners. Clark and his colleagues (Clark, 1996; Clark & Carlson, 1982) further proposed that speakers formulate utterances by consulting information that is mutually shared with their partner, or common ground. Some studies have found that speakers tailor their messages based on the listener's information needs (Lockridge & Brennan, 2002). For example, speakers tended to mention atypical (e.g., stabbed with a knife) than typical instruments (e.g., stabbed with an icepick); and introduced atypical instruments with indefinite articles (e.g., an icepick). However, others found no effect of audience design (Brown & Dell, 1987). The lack of effect of audience design (on message generation) has been attributed to speaker-internal factors (e.g., time pressure or lack of processing resources) (Arnold, 2008) or cognitive factors (e.g., saliency/prominence of information).

Goal bias is a cognitive factor that has two levels of representation, i.e., the linguistic level and the nonlinguistic level. Linguistically, goal bias refers to the bias that speakers tend to mention goal paths more frequently than source paths. Non-linguistically, goal bias refers to the bias that speakers are better at memorizing goal paths than source paths. Goal bias has been robustly found in linguistic expressions of motion events (Chen et al., 2024; Do et al., 2020, 2022) and memory/nonlinguistic representations of motion events using memory tasks (Chen et al., 2024; Do et al., 2020; Lakusta & Carey, 2015; Lakusta & DiFabrizio, 2017; Lakusta & Landau, 2012; Lakusta et al., 2007, 2017; Papafragou, 2010; see Chen et al. (2024) who found no effect of goal bias in memory where a forced-choice task was used). The combined investigation of audience design and goal bias enables us to look into how audience design might be constrained by the cognitive factor of goal bias.

Besides, language bias also affects the message generation processes (Liao et al., 2020). Slobin (1996)'s 'Thinking for Speaking hypothesis' maintains that language-specific encoding biases may affect the way speakers choose the conceptual components they will communicate, the perspective from which they want to construct the conceptual components and how they organize these conceptual components during the stage of message generation (Bock, 1995; Levelt, 1989). As Chinese differs from English in linguistic expressions of goal bias, the crosslinguistic difference in source-goal encoding prompts us to examine whether this difference may affect Chinese speakers' way of planning their message for motion events. Therefore, this study aims to explore the role of audience design and goal bias in Chinese speakers' message generation.

1.1. Audience design and goal bias that affect message generation

As mentioned above, whether speakers take into account their addressees' information needs or knowledge state during the initial message generation stage remains controversial. Several studies supported the notion that speakers take their listeners into account at the initial message generation stage (Do et al., 2020, 2022; Lockridge & Brennan, 2002; Vanlangendonck et al., 2016). In contrast, some other studies did not find the effect of audience design on message generation (Brown & Dell, 1987; Horton & Gerrig, 2002; Horton & Keysar, 1996). While numerous studies have examined the impact of audience design on reference production through classic referential communication tasks (Brennan et al., 2010; Brennan & Clark, 1996; Brown-Schmidt & Tanenhaus, 2006), research in the event domain remains comparatively limited (Brown & Dell, 1987; Do et al., 2020, 2022; Grigoroglou & Papafragou, 2019; Lockridge & Brennan, 2002). Brown and Dell (1987) instructed speakers to read and retell stories that include events with typical (e.g., stabbing with a knife) or atypical (e.g., stabbing with an icepick) instruments to addressees who could or could not have visual access to the story pictures. Their findings indicated that speakers' decisions to mention the instruments remain unaffected by the addressees' knowledge state, regardless of whether the addressees could see the events or not. However, Lockridge and Brennan (2002) conducted a modified version of Brown and Dell (1987)'s study and obtained contrasting results. They discovered that participants were more inclined to mention atypical instruments, doing so early in their retelling, and also marked atypical instruments as indefinite (e.g., an icepick) when the addressees could not see the picture. These findings suggest that speakers take addressees' belief or knowledge state into account during the early stages of message generation. Lockridge and Brennan (2002) attributed this conflicting outcome to the fact that the listener was a confederate of the experimenter, whereas in their study, the listener was a naïve participant. In a similar investigation, Grigoroglou and Papafragou (2019) manipulated the profile of the listener and discovered that both adults and children were more likely to add information about instruments when communicating with an interactive listener as opposed to a silent listener. These findings demonstrate that the involvement of the addressee and the speaker's assumptions about this involvement influence the processes of message generation (Brennan et al., 2010; Kuhlen & Brennan, 2010).

Linguistic evidence of goal bias has been robustly confirmed in experiments (Chen et al., 2024; Do et al., 2020, 2022), corpora (Arnold, 2001; Stefanowitsch & Rohde,

2004), crosslinguistic studies in Arabic, Chinese, English (Regier & Zheng, 2007) and Greek (Johanson et al., 2019) and even in special groups, e.g., children with Williams syndrome (Lakusta & Landau, 2005) and deaf children (Zheng & Goldin-Meadow, 2002). For example, Lakusta and Landau (2005) found that both typically developing children and those with Williams syndrome showed a stronger preference for expressing goal paths over source paths, even when instructed to use verbs biased toward source paths.

For the nonlinguistic evidence of goal bias, studies using memory tasks have confirmed the goal bias in speakers' conceptual representations of events. Chen et al. (2024) showed that goal bias disappears under a forced-choice task, which provides an easier retrieval context for sources to be 'reinstated'. Despite the contrary results, they proposed that the goal is robustly prominent in event representation but the source is relatively fragile, and confirmed that goal bias exists robustly in the memory task. The goal bias has been verified in infants (Lakusta et al., 2007; Lakusta et al., 2017; Lakusta & Carey, 2015; Lakusta & DiFabrizio, 2017), children (Lakusta & Landau, 2012; Papafragou, 2010) and adults (Chen et al., 2024; Do et al., 2020; Regier, 1996; Regier & Zheng, 2007) through memory tasks. For instance, Lakusta et al. (2017) revealed that infants can categorize goal paths, but not source paths, as early as 10 months. Children and adults are more accurate at detecting changes to the goal path than changes to the source path (Lakusta & Landau, 2012; Papafragou, 2010). These results indicate a cognitive-attentional or nonlinguistic bias favoring goals over sources in spatial representation and memory.

In Do et al. (2020)'s study, they explored the combined effect of audience design and goal bias on English speakers' message generation during source-goal motion event descriptions. They manipulated whether the source (in Experiment 1) or the goal (in Experiment 2) was known or unknown to an engaged addressee. Each experiment comprised a description task and a postverbal memory task. In the description task, participants verbally described the motion event observed in a video clip to their addressees. In the postverbal memory task, they detected the changes manipulated with the source or the goal in a second set of video clips. Do et al.'s findings indicated a tendency among speakers to omit the source when it was shared knowledge between speakers and addressees, compared to when it was not shared. Speakers performed less accurately in detecting source changes when the source was shared knowledge between interlocutors. Interestingly, speakers consistently mentioned goals at a high rate, regardless of whether the goals were shared between interlocutors or not. Furthermore, their detection of goal changes was consistent across conditions. Taken together, this study demonstrated that speakers consider listeners' information needs and knowledge state while generating messages for source-goal motion events, highlighting the impact of audience design. However, the fact that audience design influenced the source but not the goal suggested that its effect on English speakers' message generation for source-goal motion events was constrained by goal bias.

1.2. Goal bias in Chinese motion events

The concepts of Source and Goal are fundamental elements of a path of motion. Although the phenomenon of goal bias has been well-established in English, it has been rarely studied in Chinese. Existing evidence suggests the presence of a goal bias

in Chinese (Chen & Guo, 2009; Zheng & Goldin-Meadow, 2002; for studies of goal bias from the semantic distinction perspective, see Regier & Zheng, 2007). For example, Zheng and Goldin-Meadow (2002) analyzed how deaf and hearing children from China and the United States described motion events and revealed that both groups of children expressed more goals than sources, suggesting a goal bias in Chinese. Chen and Guo (2009) studied the manner and path encoding patterns of motion events in Chinese novels. They summarized all types of path verbs and counted the frequency of each type of path verb (p. 1757). Based on their data, we conducted a chi-square test to compare the frequency of source path verbs (e.g., 'exit') and goal path verbs (e.g., 'enter'). The result showed that the frequency of goal path verbs was significantly higher than that of source path verbs ($\chi^2(1) = 16.04$, $p = 0.000$). Because goal-path verbs are usually followed by a goal, whereas source-path verbs are usually followed by a source, these results may imply that there is a goal bias in Chinese speakers' linguistic representation of sources and goals.

Crosslinguistic differences in goal bias have been observed between Chinese and English. This difference may stem from their typological distinctions. According to Talmy (1985), languages can be classified into satellite-framed and verb-framed languages in terms of the lexicalization patterns in motion event descriptions. The core difference between these two typologies lies in whether the path of motion is encoded in the verb root. For satellite-framed languages like English, the path of motion is usually encoded in a satellite position, i.e., a post-verbal prepositional phrase like '*into the shop*' in the sentence '*The man walked into the shop*'. However, for verb-framed languages like Spanish, the path of motion is usually encoded by the main verb, like the verb *salió* 'exited' in the Spanish sentence *la botella salió de la Cueva flotando* 'the bottle exited the cave floating'. For the typology of Chinese, evidence has emerged that Chinese is not a typical satellite-framed or verb-framed language, but shares features of both in manner/path encoding and should be categorized as an equipollently framed language (see Chen & Guo, 2009; Ji & Hohenstein, 2017; Zhao & Hu, 2018). Chinese, as an equipollently-framed language, exhibits a higher frequency of path expressions compared to English (a satellite-framed language) (Zhao & Hu, 2018; Zhao & Li, 2022). That is to say, Chinese speakers may mention more goals than English speakers, as goals typically follow paths. However, this does not necessarily mean that Chinese is as strongly goal-biased as English. The reason is simple: to determine whether a language is goal-biased, one should compare the frequency of goal mentions to that of source mentions. Liu and Wen (2023) investigated how Chinese and English speakers encode motion events. The results revealed that while English native speakers mentioned more goals ($M = 0.036$, $SD = 0.010$) than sources ($M = 0.011$, $SD = 0.010$), Chinese native speakers' goal mentions ($M = 0.039$, $SD = 0.010$) were slightly higher than source mentions ($M = 0.036$, $SD = 0.011$). This asymmetry may lead to a potentially weaker goal bias in Chinese than in English.

Based on the above evidence, we hypothesize that Chinese exhibits a goal bias, but this bias is relatively fragile in comparison to the strong goal bias observed in English (we will revisit this issue in the General Discussions part). However, as mentioned in the introduction, two criteria have been identified to determine whether a language exhibits goal bias: linguistic and nonlinguistic criteria. Linguistically, goal bias is indicated by speakers mentioning goals more frequently than sources. Nonlinguistically, it is demonstrated by speakers having a better memory for goals than for sources. Our hypothesis above is based on some indirect evidence. Therefore, more

linguistic and nonlinguistic evidence is needed to test the hypothesis. This is the second aim of the present study.

While previous studies have highlighted differences in message generation of manner and path among speakers of typologically different languages (Ji & Hohenstein, 2017; Zhao & Hu, 2018), less attention has been paid to the differences in message generation of path elements, with current research limited to Trajectory and Goal (Flecken et al., 2015; Liao et al., 2020). Liao et al. (2020) examined how speakers of Chinese and Dutch conceptualized the path of motion differently. However, their study focused solely on the path elements like Trajectory and Goal, but not Source. The question of whether the language-specific encoding of goal bias in Chinese influences the way Chinese speakers generate messages about source-goal motion events remains unanswered.

1.3. *The present study*

This research was conducted to investigate the effects of audience design and goal bias on Chinese speakers' message generation. The reasons to choose these two factors are as follows. First, as for the factor of audience design, on the one hand, there exists a debate on whether speakers consider their listeners' knowledge state during the message generation stage. While Do et al. (2020) reported the effects of audience design on English native speakers' message generation, Brown and Dell (1987) did not find such effects. On the other hand, audience design is indeed a pragmatic factor. It is pragmatic because it is grounded in the reality of how people actually behave and what they actually need. The effects of audience design observed in one language group may not necessarily be observed in another language group. Therefore, whether audience design affects Chinese native speakers' message generation deserves to be investigated. Second, whether the language-specific goal bias in Chinese influences Chinese speakers' message generation of source-goal motion events has not been touched upon yet. According to Slobin (1996)'s 'Thinking for Speaking' hypothesis, language-specific encoding biases may affect the way speakers choose the conceptual components they will communicate, the perspective from which they want to construct the conceptual components and how they organize these conceptual components during the stage of message generation (Bock, 1995; Levelt, 1989). Therefore, whether the fragile goal bias in Chinese influences Chinese speakers' message generation in the same way as it does for English speakers is an interesting question. Based on the above reasons, our research aims to answer this question: If and to what extent do audience design and goal bias influence Chinese speakers' message generation of source-goal motion events? If the answer is 'yes', then what characteristics manifest in the message generation of source and goal among Chinese speakers?

This study investigated Chinese speakers' message generation of source-goal motion events through a description task and a memory task. The reasons lie in the fact that goal bias has been defined as a cognitive factor with both linguistic and nonlinguistic representations, and previous studies have used both description tasks and memory tasks to examine goal bias at these two levels (Do et al., 2020; Papafragou, 2010). The description task tests which event elements speakers choose to say and which they choose not to say. The memory task tests speakers' accuracy in memorizing sources and goals after the immediate description of motion events.

The pragmatic conditions were manipulated as follows: the status of the source (in Experiment 1) or the goal (in Experiment 2) was manipulated as known (common ground) or unknown (no common ground) to conversationally engaged addressees, creating a natural communicative setting. This approach, which features an engaged addressee rather than an experimenter or an imagined pseudo-listener, was adopted to enhance the ecological validity of this study.

2. Experiment 1

2.1. Purpose

Experiment 1 explored if and to what extent audience design and goal bias influence Chinese speakers' message generation of source-goal motion events by manipulating sources as known or unknown to the confederate addressees.

2.2. Design

Experiment 1 adopted a language description task and a memory task (Do et al., 2020; Lakusta et al., 2007). The description task required participants to describe the motion clip to their copresent addressee with the source shown (common ground) or not shown to the addressee (no common ground). The follow-up memory task required participants to decide whether the second set of video clips was identical to the one they described in the description task. The second set of video clips may involve: (1) changing the source; (2) changing the goal or (3) no change at all. Clips in the memory task were presented in the same order as in the description task. The memory task was unexpected and administered once participants had finished the description task.

The description task employed a 2 (mention type: source versus goal) $\times$ 2 (ground type: common ground condition versus no common ground condition) mixed design. Mention type served as a within-subject factor, and ground type as a between-subject factor. Similarly, the memory task utilized a 3 (change type: source change versus goal change versus no change) $\times$ 2 (ground type: common ground condition versus no common ground condition) mixed design. Change type served as a within-subject factor, and ground type as a between-subject factor.

2.3. Participants

Sixty-eight participants from a university in South China (aged 19–26; 34 females and 34 males) were recruited for the experiment. The required number of participants was determined using G*Power based on a power analysis of the effect sizes reported in Experiment 1 by Do et al. (2020). The analysis indicated that, with $\alpha = 0.05$ and $\beta = 0.80$, the projected sample size was $n = 37$, when odds ratio = 26.05 (calculated from the estimate value $\beta = 3.26$). Due to differences in the number of experimental items (24 in our study and 18 in Do et al., 2020, at least 28 participants per ground-type condition were needed. Finally, 34 participants per ground-type condition were recruited, resulting in a total of 68 participants.

Besides, two confederate addressees were employed, with each assigned to one of the two ground-type conditions. Participants were intentionally led to believe that

these confederate addressees were fellow participants. After the experiment, all participants received a gift as a reward.

2.4. Materials

The target source-goal motion clips depicted an animate figure moving from an inanimate source object (i.e., the starting point) to an inanimate goal object (i.e., the endpoint). The preparation of these target motion clips involved identifying the figure, source, goal items and the manner verbs comprising these clips. The manner verbs were selected from verbs frequently used in Chinese novels. Specifically, the manner verbs in this study included ‘飞 *fei*¹ “fly”’, ‘爬 *pa*² “crawl or climb”’ and ‘跳 *tiao*⁴ “jump”’. Once the manner verbs were selected, specific figure, source and goal items were configured. Finally, 24 target motion clips were created. To assess the acceptability and reasonableness of these clips, 20 Chinese speakers rated the 24 motion clips on a 1–5 Likert scale questionnaire (1 indicating completely unreasonable, 5 indicating completely reasonable). Motion events scoring above 3 were selected for the target motion clips. 24 target motion clips were finalized.

Speakers could view the entire motion clip in both ground-type conditions. For addressees, they could only view the starting point of a source-goal motion event in the common ground condition, but could not view any part of the motion event in the no common ground condition. Figure 1 demonstrates an example of a motion clip from the speaker’s view on a source-goal motion event, including the beginning (source), middle (medium) and the endpoint (goal) of a source-goal motion event. Figure 2 presents an example of a motion clip from the addressee’s view, which includes only the starting point of a source-goal motion event.

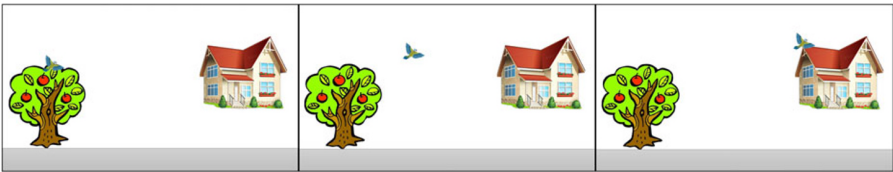


Figure 1. Speaker’s view of a motion clip in both ground-type conditions.

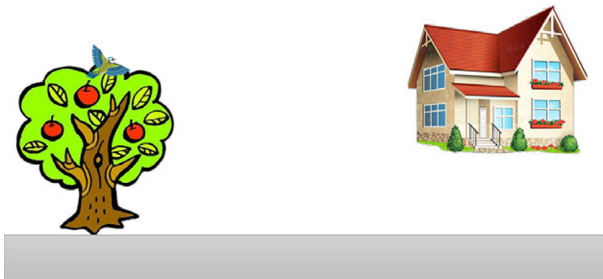


Figure 2. Addressee’s view of a motion clip in the common ground condition.

The 24 target motion clips were divided into two lists. The location of source items and goal items, as well as the path direction, were counterbalanced. Items functioning as sources in one list were balanced by serving as goals in another list. The path direction in both lists was counterbalanced, with the figure moving from left to right in half of the motion clips and from right to left in the other half. 24 filler motion clips were created that did not involve either the source or the goal (see Figure 3 below).

For the memory task, the same 24 target motion clips in the description task were used and manipulated as follows: 8 motion clips with the source changed, 8 with the goal changed and the remaining 8 left unchanged. Following Papafragou (2010)'s approach, changes to both source and goal items involved substitutions with items from the same category (see Figure 4 for a source-changed motion clip).

2.5. Procedures

Participants watched brief video clips and then described them to the confederate addressee (34 participants with a female addressee and 34 with a male addressee). They were led to believe that their partner would answer a simple question about the video clip on a separate screen based on their descriptions, and also led to believe that their partners, who acted as confederate addressees in reality, were also subjects under investigation. To control for engagement, confederate addressees were required to show similar levels of involvement with those speakers by keeping eye contact or using simple verbal responses like 'yes' or 'no' to signal readiness for the next trial across two ground-type conditions. Participants performed two practice trials before moving to the formal experiment.

In the common ground condition, speakers and addressees initially sat side by side. Both the speaker and the addressee viewed the motion clip's first frame, the source (the starting point) on the speaker's screen (see Figure 2). Then, after viewing the first frame, the addressee returned to the seat across from the speaker and waited for the speaker's description of the video clip. In the no common ground condition,

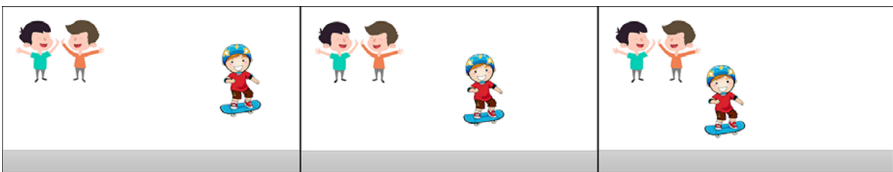


Figure 3. A motion clip of distractor displays.



Figure 4. A source-changed sample motion clip.

speakers and the addressees sat across from each other during the whole process, ensuring that the addressees could not see any part of the video clips. The speakers needed to describe each motion clip to their copresent addressee. Their utterances were recorded.

The memory task was unexpected and administered once participants had finished the description task. They were instructed to watch a second set of video clips and to judge whether these were the same as or different from those viewed in the description task. Participants marked the *same* column on the answer sheet if the motion clip was the same as the one in the previous description task, and a *different* column if it was not. They had no time limit and were solely tested on their accuracy in detecting source or goal changes.

2.6. Coding

A total of 1440 utterances were transcribed. We coded whether participants mentioned source/goal or not (mention coded with 1, no mention coded with 0). All mentions within the following syntactic frames were included: (1) preposition + source (从树上 'from the tree'); (2) manner verb + path verb + goal (爬到鞋子上 'crawl to the shoes', 飞到台灯 'fly to the lamp', 跳上床 'jump onto the bed'); (3) path verb + goal (上了滑梯 'onto the ladder'); (4) source + location + figure (铁塔上的鹰 'the eagle on the towel').

2.7. Results

2.7.1. Linguistic mentions of sources and goals

Figure 5 illustrates the proportion of source and goal mentions in both the common ground and no common ground conditions. In the common ground condition, the participants mentioned sources in 75.3% of utterances ($SD = 0.42$) and goals in all utterances ($SD = 0$). In the no common ground condition, the participants mentioned sources in 91.7% of utterances ($SD = 0.10$) and goals in all utterances ($SD = 0$). The participants exhibited a higher frequency of mentioning sources and goals in the no common ground condition compared to the common ground condition.

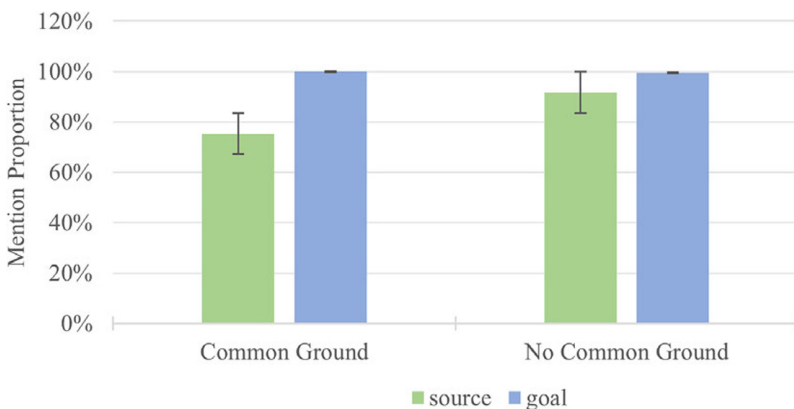


Figure 5. Proportion of source and goal mentions in the common and no common ground conditions.

Statistical analyses of the proportion of source and goal mentions were conducted using the generalized logistic mixed-effect model through the lme4 package in R (R Development Core Team, 2018). Fixed-effect factors included ground type (common ground versus no common ground) and mention type (source versus goal), both sum-coded. The dependent variable is whether participants mentioned the source or the goal. Mention type was included as part of the by-participant and by-item random effects; Ground type was only included as part of the by-item random effects. The maximal random effect structure was maintained (Barr et al., 2013; Brown, 2021) and the model was simplified only if it failed to converge. The final model that converged included ground type and mention type as fixed effects, by-participant and by-item random intercepts as well as by-participant random slopes of mention type and by-item random slopes of the interaction between ground type and mention type.

Table 1 displays the estimated fixed effects from logistic mixed-effects models for source and goal mentions in the description task. The main effects of ground type and mention type were neither significant ($p > 0.05$). The interaction between ground type and mention type was not significant ($p > 0.05$). As the goal mentions in both common ground and no common ground conditions reached 100% mentions, the dependent variable lacks variance in this subset of the data. Since logistic regression requires variability in the binary outcome (i.e., values distributed across both 0 and 1), the model cannot be estimated under such conditions. Therefore, from a statistical standpoint, it is not feasible to assess the effect of ground type in the goal-mention condition. As a result, we focus our analysis on the source mention condition. The results showed that Chinese speakers mentioned more sources in the common ground condition than in the no common ground condition ($\beta = -2.130$, $SE = 0.769$, $z = -2.770$, $p = 0.006$).

2.7.2. Accuracy rate of change detection for sources and goals

Figure 6 summarizes the correct detection rates of the source change, the goal change and no change condition. In the common ground condition, the accuracy rates for detecting the source change, the goal change and no change were 0.628 ($SD = 0.48$), 0.807 ($SD = 0.39$) and 0.88 ($SD = 0.12$), respectively. In the no common ground condition, the accuracy rates were 0.742 ($SD = 0.43$) for the source change, 0.829 ($SD = 0.365$) for the goal change and 0.906 ($SD = 0.10$) for the condition of no change. Overall, the participants demonstrated higher accuracy in the no common ground condition than in the common ground condition.

Statistical analyses were conducted on the accuracy data using a generalized logistic mixed-effect model. Ground type (common ground vs no common ground)

Table 1. Generalized logistic mixed-effects model (GLMEM) estimates of fixed effects for source and goal mentions in the description task

	Estimate	S.E.	Z	p
Intercept	13.634	14.777	0.923	0.356
Mention type	20.925	29.548	0.708	0.479
Ground type	4.895	29.327	0.167	0.867
Mention type $\times$ Ground type	14.073	58.651	0.240	0.810

Note: The full model syntax was: `glmer(Mention ~ MentionType*GroundType + (1 + MentionType|ParID) + (1 + GroundType*MentionType|itemID))`.

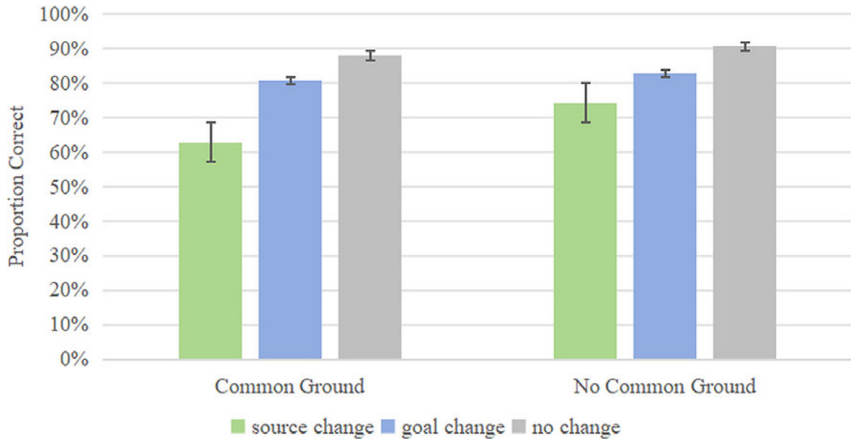


Figure 6. Accuracy rate for the source change, the goal change and no change conditions in the common ground and no common ground conditions.

and change type (source change versus goal change) were included as fixed effects, with participants and items as random effects. The final model that converged included ground type and change type as fixed effects, by-participant and by-item random intercepts, as well as by-participant random slopes of change type and by-item random slopes of the interaction between ground type and change type. The no change condition was not included in the analysis as it served merely as a baseline accuracy indicator, showing similar results in both the common ground ($M = 0.88$, $SD = 0.12$) and no common ground condition ($M = 0.906$, $SD = 0.10$).

Table 2 displays the estimated fixed effects from generalized logistic mixed-effects models for the correct detection rates for the source change and the goal change in the memory task. The main effect of ground type was not significant ($\beta = -0.467$, $SE = 0.293$, $z = -1.591$, $p = 0.112$). The main effect of change type was marginally significant ($\beta = 0.773$, $SE = 0.415$, $z = 1.862$, $p = 0.063$), suggesting that participants showed a numerical tendency to detect goal changes more accurately than source changes. There was no interaction effect between ground type and change type ($p > 0.05$). To explore if the main effect of change type was different across ground-type conditions, separate models were built for the two levels of ground type. We used the Benjamini–Hochberg procedure to correct for multiple comparisons. Results showed that participants were better at detecting the goal change than the source change in the common ground condition ($\beta = 1.271$, $SE = 0.504$, $z = 2.523$, $p = 0.024$, *FDR-corrected*), but not in the no common ground condition ($p > 0.05$).

Table 2. Generalized logistic mixed-effects model (GLMEM) estimates of fixed effects for the correct detection rates for the source change and the goal change in the memory task

	Estimate	S.E.	z	p
Intercept	1.652	0.324	5.093	< 0.001
Change type	0.773	0.415	1.862	0.063
Ground type	−0.467	0.293	−1.591	0.112
Change type × Ground type	0.607	0.528	1.151	0.250

Note: The full model syntax was: `glmer(Accuracy~ ChangeType*GroundType + (1 + ChangeType|ParID) + (1 + ChangeType*GroundType|itemID))`.

2.8. Discussion

In the description task, the results revealed that speakers were more likely to omit the source in motion event descriptions when the source information was shared with addressees, compared to when it was not, highlighting the impact of audience design on source mentions. This result aligned with previous findings in English, where speakers adhered to Grice's (1975) pragmatic principle Maxim of Quantity, aiming to be informative while being concise in language communication (Do et al., 2020, 2022; Grigoroglou & Papafragou, 2019).

In the memory task, participants demonstrated similar accuracy in detecting source and goal changes across common ground and no common ground conditions, demonstrating an absence of effect of audience design on speakers' memory of sources. This result contradicts the findings of Do et al. (2020), who found that audience design can modulate English speakers' memory of sources. We discussed the reasons underlying these divergent findings in the General Discussion section.

For the goal bias pattern at the linguistic level, participants mentioned sources and goals at a similar rate, indicating no goal bias in Chinese speakers' linguistic representation of sources and goals. For the goal bias pattern at the nonlinguistic/memory level, participants were better at detecting the goal change than the source change in the common ground condition, but not in the no common ground condition, indicating a fragile goal bias in their memory representation of sources and goals.

3. Experiment 2

3.1. Purpose

Experiment 2 explored if and to what extent audience design and goal bias influence Chinese speakers' message generation of source-goal motion events by manipulating goals, whether known or unknown to the confederate addressees.

3.2. Design

The design of Experiment 2 was similar to Experiment 1. The only difference is that the goal, instead of the source, was manipulated as shown (common ground) or not shown (no common ground) to the addressee in the description task.

3.3. Participants

An additional sixty-eight Chinese speakers (aged 19–26; 34 females and 34 males) from a university in South China participated in Experiment 2. They received a gift as a reward after the experiment.

3.4. Materials

The materials were identical to those in Experiment 1, with two exceptions. First, participants viewed the last frame of the video clip (the endpoint of a motion event). Note that addressees can view it only in the common ground condition. Second, each motion clip was preceded by a 'reply' screen (see Figure 7) to inform participants

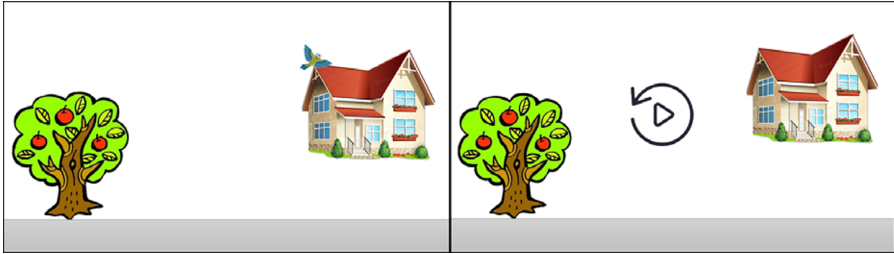


Figure 7. Both speaker and addressee view the last frame in common ground.

that they were viewing the last frame of the motion clip. Figure 7 illustrates the shared view of the last frame between the speaker and the addressee in the common ground condition. After viewing Figure 7, the speaker watched an unfolding motion clip identical to the one in Figure 1. The addressee, however, could not see the unfolding portion depicted in Figure 1. In the no common ground condition, the addressee cannot view any part of the motion clip.

3.5. Procedures

Participants described the motion clip to the addressee and conducted a follow-up memory task. The procedures in Experiment 2 were identical to those in Experiment 1 with the following exceptions. In the description task, firstly, the last frame of the video clip was presented. In the common ground condition, both speakers and addressees viewed the last frame of the video clip. Participants were instructed to watch a replay of this video clip and describe it to their partners. After viewing the last frame, the addressee sat across from the speakers. Subsequently, they were required to click the mouse to proceed to the 'begin replay' screen. Only the speakers could watch the replay of the video clip. In the no common ground condition, the addressees were unable to see any part of the motion clips.

3.6. Coding

Data coding in Experiment 2 followed the same methodology as in Experiment 1.

3.7. Results

3.7.1. Linguistic mentions of sources and goals

Figure 8 displays the proportion of source and goal mentions in both the common ground and no common ground conditions. In the common ground condition, speakers mentioned sources in 96.3% of the utterances ($SD = 0.20$) and goals in 99.6% of the utterances ($SD = 0.06$). In the no common ground condition, speakers mentioned sources in 96.2% of the utterances ($SD = 0.20$) and goals in 99.8% of the utterances ($SD = 0.04$).

Statistical analyses of the proportion of source and goal mentions were conducted using the same criteria as those in Experiment 1. The final model that converged included ground type and mention type as fixed effects, by-participant and by-item random intercepts as well as by-participant random slopes of mention type and

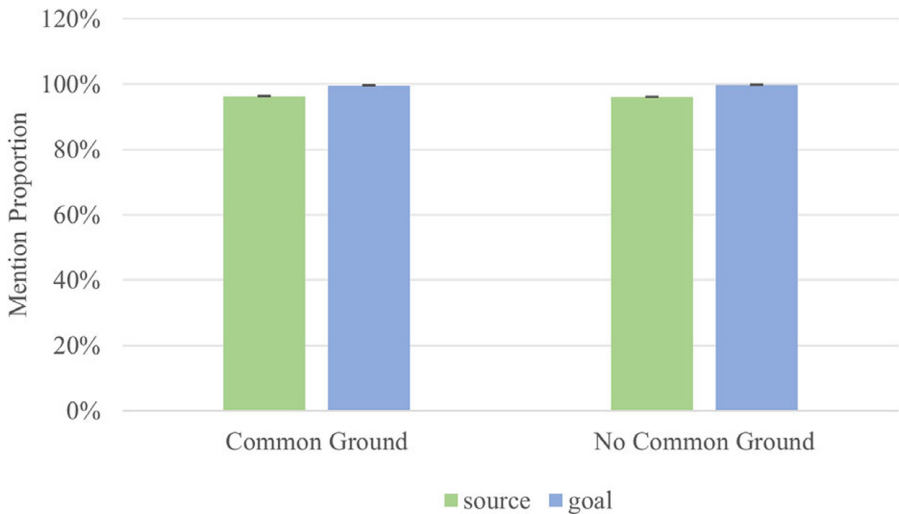


Figure 8. Proportion of source and goal mentions in the common and no common ground conditions.

Table 3. Generalized logistic mixed-effects model (GLMEM) estimates of fixed effects for source and goal mentions in the description task

	Estimate	S.E.	z	p
Intercept	10.571	1.792	5.898	< 0.001
Mention type	3.348	3.268	1.024	0.306
Ground type	1.138	1.649	0.690	0.490
Mention type × Ground type	−0.050	2.317	−0.022	0.983

Note: The full model syntax was: `glmer(Mention~ MentionType*GroundType+ (1 + MentionType|ParlD) + (1 + MentionType +GroundType|itemID))`.

by-item random slopes of ground type and mention type. Table 3 showed that the main effect of ground type was not significant ($p > 0.05$), suggesting no notable differences in the frequency of source and goal mentions between the two ground-type conditions. The main effect of mention type was neither significant ($p > 0.05$), indicating a similar frequency of source mentions and goal mentions across the two ground-type conditions. There was no interactive effect between ground type and mention type ($p > 0.05$).

3.7.2. Accuracy rate of change detection for sources and goals

Figure 9 summarizes the correct detection rates of the source change, the goal change and no change condition. In the common ground condition, the accuracy rates for detecting the source change, the goal change and the no change were 0.742 ($SD = 0.43$), 0.738 ($SD = 0.45$) and 0.935 ($SD = 0.08$), respectively. In the no common ground condition, the accuracy rates for detecting source change, goal change and no change were 0.753 ($SD = 0.43$), 0.801 ($SD = 0.41$) and 0.919 ($SD = 0.09$), respectively.

Statistical analyses were conducted by using the same criteria as those in Experiment 1. The no change condition was not included in the analysis as it served merely as a baseline accuracy indicator, showing similar results in both the common ground

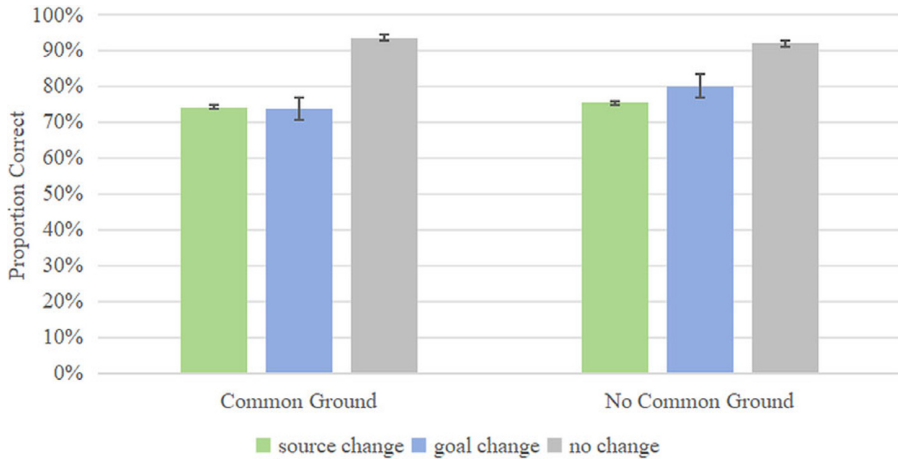


Figure 9. Accuracy rate for the source change, the goal change and no change conditions in the common ground and no common ground conditions.

Table 4. Generalized logistic mixed-effects model (GLMEM) estimates of fixed effects for the correct detection rates for the source change and the goal change in the memory task

	Estimate	S.E.	Z	p
Intercept	1.513	0.281	5.385	< 0.001
Change type	−0.190	0.578	−0.328	0.743
Ground type	−0.118	0.254	−0.463	0.644
Change type × Ground type	−0.378	0.575	−0.658	0.511

Note: The full model syntax was: `glmer(Accuracy ~ ChangeType*GroundType + (1 + ChangeType|ParID) + (1 + ChangeType*GroundType|itemID))`.

($M = 0.935$, $SD = 0.08$) and no common ground condition ($M = 0.919$, $SD = 0.09$). The final model that converged included ground type and change type as fixed effects, by-participant and by-item random intercepts, as well as by-participant random slopes of change type and by-item random slopes of the interaction between ground type and change type. Table 4 shows that the main effect of ground type was not significant ($p > 0.05$), suggesting no notable differences in the detection of source and goal changes between the two ground-type conditions. Similarly, the main effect of change type was not significant ($p > 0.05$), indicating no significant difference in detecting source changes compared to goal changes across both ground-type conditions. Additionally, no significant interaction effect was found for ground type and change type ($p > 0.05$).

3.7.3. Joint data analysis

To determine if audience design impacts sources and goals similarly, we compared the mentions and accuracy rates for change detection of sources and goals across Experiment 1 and Experiment 2. Both analyses were conducted using the logistic mixed-effect model. The final model with the maximal structure that converged for the description task included mention type (source versus goal), ground type

(common ground versus no common ground) and experiment (1 versus 2) as fixed-effect factors, by-participant and by-item random intercepts as well as by-participant random slopes of mention type and by-item random slopes of experiment and the interaction between ground type and mention type. The final model with the maximal structure that converged for the memory task included change type (source change versus goal change), ground type (common ground versus no common ground) and experiment (1 versus 2) as fixed-effect factors, by-participant and by-item random intercepts as well as by-participant random slopes of change type and by-item random slopes of experiment and the interaction between ground type and change type.

Table 5 shows the generalized logistic mixed-effects model estimates of fixed effects for source and goal mentions in the description task across two experiments. The main effect of mention type reached significance ($\beta = 10.018$, $SE = 4.723$, $z = 2.121$, $p = 0.034$), showing that participants mentioned more goals than sources across the two experiments. There was no effect of ground type or experiment, nor were there any interactions between mention type, ground type and experiment ($p > 0.05$).

Table 6 presents the generalized logistic mixed-effects model estimates of fixed effects for the accuracy rates in detecting source and goal changes across the two experiments. No main effects of change type, ground type and experiment were found ($p > 0.05$). The interaction between experiment and change type was significant ($\beta = -0.786$, $SE = 0.388$, $z = -2.029$, $p = 0.043$). Post hoc analyses revealed that

Table 5. Generalized logistic mixed-effects model (GLMEM) estimates of fixed effects for source and goal mentions in the description task across two experiments

	Estimate	S.E.	z	p
Intercept	10.280	2.381	4.317	< 0.001
Experiment	-1.484	4.322	-0.343	0.731
Mention type	10.018	4.723	2.121	0.034
Ground type	3.603	4.323	0.833	0.405
Experiment × Mention type	-10.431	8.555	-1.219	0.223
Experiment × Ground type	-5.869	8.489	-0.691	0.489
Mention type × Ground type	7.987	8.600	0.929	0.353
Experiment × Mention type × Ground type	-19.779	16.866	-1.173	0.241

Note: The full model syntax was: `glmer(Mention ~ Experiment* MentionType*GroundType + (1 + MentionType | ParID) + (1 + MentionType*GroundType + Experiment|itemID))`.

Table 6. Generalized logistic mixed-effects model (GLMEM) estimates of fixed effects for the accuracy rate in the memory task across two experiments

	Estimate	S.E.	z	p
Intercept	1.601	0.284	5.627	< 0.001
Experiment	0.021	0.208	0.099	0.921
Change type	0.248	0.318	0.779	0.436
Ground type	-0.267	0.193	-1.382	0.167
Experiment × Change type	-0.786	0.388	-2.029	0.043
Experiment × Ground type	0.237	0.345	0.685	0.493
Change type × Ground type	0.110	0.367	0.300	0.764
Experiment × Change type × Ground type	-0.978	0.690	-1.417	0.156

Note: The full model syntax was: `glmer(Accuracy ~ Experiment*ChangeType*GroundType + (1 + ChangeType | ParID) + (1 + GroundType*ChangeType*Experiment|itemID))`.

the main effect of change type was marginally significant in Experiment 1 ($\beta = -0.641$, $SE = 0.373$, $z = -1.719$, $p = 0.086$), confirming the finding of a numerical trend for participants to detect goal changes more accurately than source changes in Experiment 1. However, this effect was not observed in Experiment 2 ($p > 0.05$). Additionally, the main effect of experiment was not significant for either the source change condition or the goal change condition ($p > 0.05$).

In summary, whether ground type affects the mention or memory of sources and goals, the results revealed that participants mentioned more sources in the no common ground condition than in the common ground condition. However, they showed similar accuracy in detecting the source change across the ground type conditions, suggesting the effect of ground type on source mention but not on source memory. On the contrary, participants mentioned similar goals and demonstrated similar accuracy in detecting the goal change across two ground-type conditions, suggesting the absence of effect of ground type on goal mention and goal memory.

Regarding source/goal mentions, participants mentioned more goals than sources in the combined analysis of both experiments, but this pattern was not observed in the individual experiment. This suggests a fragile rather than robust goal bias in Chinese speakers' linguistic representation of motion events. With respect to source/goal memory, goal bias was only observed in the common ground condition in Experiment 1, but not in other experimental conditions, indicating a fragile goal bias in memory representation of motion events.

3.8. Discussion

Even when goals were known to the addressees and over-informative for communication, speakers still described them to listeners, indicating that audience design did not influence goal mentions. This result contrasted with Experiment 1, where source mentions decreased when known to addressees, yet aligned with findings in English studies (Do et al., 2020). The analysis comparing Experiments 1 and 2 confirmed that audience design influenced speakers' mentions of sources but not of goals. The lack of audience design effect on goal mentions may stem from the high prominence of goals in event representation. Similar to English, goals in Chinese are also salient in event representation (Lakusta et al., 2007; Lakusta & Carey, 2015; Lakusta & Landau, 2005; Papafragou, 2010; Regier & Zheng, 2007). Another reason was that for completed motion events tested in this study, speakers may use the goal as a linguistic device to convey these motion events, even if this information was shared with the addressees.

Speakers showed similar accuracy in detecting goal changes regardless of goals known or unknown to addressees, suggesting that the speakers' memory for goals was not influenced by the listeners' knowledge state. This finding reaffirmed the high prominence of goals in event representation.

For the goal bias pattern at the linguistic level, it was found only in the combined analysis of two experiments, suggesting a fragile goal bias pattern in Chinese speakers' linguistic representation of motion events rather than a robust goal bias pattern found in native English speakers (Chen et al., 2024; Do et al., 2020). For the goal bias pattern at the nonlinguistic/memory level, goal bias existed only in the common ground condition in Experiment 1 but disappeared across other experimental conditions. This also revealed a fragile goal bias pattern in Chinese speakers' memory representation of sources and goals, in comparison to the robust goal bias consistently

found in English using the memory task (Chen et al., 2024; Lakusta et al., 2007; Papafragou, 2010).

4. General discussions

This study aimed to explore whether and to what extent audience design and goal bias influence Chinese speakers' message generation of source-goal motion events. Two experiments were conducted to manipulate the pragmatic status of either the source (Experiment 1) or the goal (Experiment 2) as known or unknown to the confederate addressee. In the description task, participants mentioned fewer sources and goals when the source was known to the addressee compared to when it was unknown. However, even when the goal was known to the addressee, participants continued to mention it at a high rate compared to the unknown condition. In the memory task, participants demonstrated similar accuracy in recalling the source regardless of whether it was known or unknown to the addressee. No significant differences were observed in their memory for the goal, regardless of whether it was known or unknown to the addressee. These results suggest that audience design affected the mention of the source, but not its memory, nor the mention or memory of the goal. Furthermore, a fragile goal bias was detected in Chinese, both linguistically and in memory. In the description task, it was observed only in the combined analysis of the two experiments. In the memory task, the goal bias was observed only in the common ground condition in Experiment 1, but not in other experimental conditions.

4.1. Audience design and goal bias affect message generation of source-goal motion events

To generate a message of an event, speakers need to build a conceptual representation of that event based on its components and the relations among these components, then encode the conceptual content they prefer to express in a specific language. Based on the pragmatic account, speakers were expected to mention the source or goal when it is unknown and thus under-informative to the addressees, and to omit source or goal information when it is already known and thus over-informative to the addressees in communication. Our findings verified that audience design, as a pragmatic factor, influenced speakers' mention of the source but did not affect their memory of it. Moreover, it had no impact on either the mention or the memory of the goal. Speakers considered listeners' knowledge state during the message generation stage of language production. They selectively omitted the source when it was shared knowledge between interlocutors. However, even the goal was shared, Chinese speakers consistently mentioned the goal at a high rate (99.7%) in source-goal motion event descriptions.

The finding that audience design influenced Chinese speakers' message generation of sources aligned with previous studies (Do et al., 2020, 2022; Vanlangendonck et al., 2016). In line with Grice's Maxim of Quantity (Grice, 1975), speakers adhered to the principle of maximal informativeness and efficiency in communication and adapted to listeners' information needs and knowledge state by omitting the source when it was already known to the addressees. The observed effect of audience design may stem from this study involving an engaged addressee who provided feedback through 'yes' or 'no' responses or keeping eye contact with speakers. The involvement

of the addressee likely prompts speakers to consider their information needs while generating messages (Brennan et al., 2010; Kuhlen & Brennan, 2010).

Audience design did not influence Chinese speakers' memory of the source, which contradicts the findings of Do et al. (2020), who reported an effect of audience design on source memory. This discrepancy may be attributed to differences in the cognitive salience of sources across languages: in Chinese, sources may occupy a more prominent position in memory representations of motion events, whereas in English, they hold a more peripheral role (Do et al., 2020, 2022; Lakusta & DiFabrizio, 2017; Lakusta & Landau, 2005). This prominence is further supported by the fragile goal bias observed in Chinese speakers' memory, where goal bias appeared only in the common ground condition of Experiment 1 and was absent in other experimental conditions.

Audience design did not influence Chinese speakers' mention or memory of the goal, in line with (Do et al., 2020, 2022). There are several possibilities. Firstly, the significant prominence of the goal in humans' conceptual event representation might lead to its over-informative mention (Lakusta et al., 2007; Lakusta & Landau, 2005; Papafragou, 2010). The goal holds a prominent position as a core component in event representation. Speakers might prefer to encode and communicate this core component in preverbal messages, even if shared with addressees. This interpretation aligned with our discovery of a goal bias pattern in the memory task, suggesting a nonlinguistic goal bias in Chinese event representation. Secondly, the motion events in our video clips were completed motion events. In completed motion events, the action reaches completion or an endpoint, making the goal central to linguistic expression. Therefore, speakers might employ the goal as a necessary linguistic element to convey completed motion events, even if this information is shared with addressees. The third possibility could stem from the animated nature of the motion clips used in the experiments. Participants described animated motion clips instead of static pictures. The animated nature of these clips might enhance the visual perception of movement and, consequently, the perception of sources and goals in the event (Zhang, 2021).

This study proved a language-general phenomenon: audience design and goal bias jointly contribute to speakers' message generation of source-goal motion events. Specifically, audience design influences speakers' message generation. However, it is constrained by the cognitive factor goal bias, i.e., it can modulate speakers' message generation of sources but not goals.

4.2. *The characteristics of goal bias in Chinese*

In the description task, goal bias was only observed in the combined analysis of the two experiments, not in Experiment 1 or Experiment 2 individually. In the memory task, it was only found in the common ground condition of Experiment 1, not in any other conditions across the two experiments. These results suggest that goal bias exists in Chinese, at least under certain circumstances. For example, when sources were shared between interlocutors, participants mentioned more goals than sources and showed better performance in detecting goal changes than source changes. This aligns with previous studies (Zheng & Goldin-Meadow, 2002).

Importantly, the goal bias has not been detected consistently in all linguistic and nonlinguistic tasks, indicating a fragile goal bias in Chinese speakers' linguistic and

nonlinguistic/memory representations of motion events. The fragile goal bias in Chinese can be attributed to the following reasons. At the nonlinguistic level, sources and goals may occupy a prominent position in speakers' memory representation of motion events. This prominence encourages speakers to include both elements as conceptual components in the message generation of source-goal motion events. Our memory data supported this view, showing no significant differences in speakers' memory of the source and the goal across three experimental conditions (specifically, the NCG condition in Experiment 1 and both CG and NCG conditions in Experiment 2). The prominence of sources and goals has also been demonstrated by Liu and Wen (2023). Using eye-tracking techniques, they showed that Chinese speakers tend to pay more attention to sources and goals compared to native English speakers. Specifically, Chinese speakers allocate greater attention to sources and goals, while English speakers focus more on trajectories. It should be noted that the memory task in this experiment was conducted immediately after the description task. Therefore, the memory representation of the source/goal may have been influenced by the linguistic expressions of sources and goals in the description task. Future research could investigate the memory representation of sources and goals in Chinese speakers while eliminating the influence of linguistic expressions.

The fragile goal bias observed at the linguistic level may stem from underlying nonlinguistic representations. Lakusta and Landau (2012) suggest that the linguistic goal bias found in English arises from nonlinguistic goal bias (restricted to animate, goal-directed motion events). Similarly, the fragile goal bias found in Chinese at the linguistic level in the present study may also result from the prominence of both sources and goals in nonlinguistic representations. The prominence of the source and the goal may lead Chinese speakers to mention both the source and the goal in most cases. However, whether nonlinguistic representations of sources and goals lead to differences in linguistic encoding needs to be proved in future research.

The fragility of goal bias in Chinese contrasted with the robust goal bias pattern in English. In English, sources are a peripheral element, whereas goals are a conceptually privileged element (Chen et al., 2024; Do et al., 2020). In English, goal bias was robust even when the source was made perceptually salient to speakers (Lakusta & DiFabrizio, 2017), and has been established in linguistic and memory tasks (at least in the memory task) (see Chen et al., 2024 who found no goal bias using a force-choice task) (Lakusta et al., 2007, 2016; Lakusta & Carey, 2015; Lakusta & DiFabrizio, 2017; Lakusta & Landau, 2005, 2012; Pace et al., 2020; Regier & Zheng, 2007). However, in Chinese, sources were mentioned as frequently as goals and were identified with comparable accuracy in most cases across two experiments. There are two main reasons for the difference in goal bias between English and Chinese. First, at the nonlinguistic level, both source and goal representations are more prominent in Chinese. In particular, the source is cognitively prominent in Chinese, whereas it is a peripheral element in English (Do et al., 2020, 2022). Second, at the linguistic level, Chinese speakers tended to express locations when describing motion events. For example, Liao et al. (2020) showed that Chinese speakers conceptualized endpoint-oriented motion by more frequently expressing locations (e.g., on the street), whereas Dutch speakers (from a satellite-framed language, like English) focused more on the trajectory.

In summary, while previous studies have been focused on message generation of manner and path among speakers of typologically different languages (Ji & Hohenstein, 2017; Zhao & Hu, 2018), our study revealed new evidence of the differences between Chinese and English in motion events in message generation of path

elements like the source and goal. Specifically, in selecting which event element (source or goal) to mention during speech planning, Chinese speakers display a fragile goal bias pattern, contrasting with the robust goal bias in English. By examining the way Chinese speakers generate messages with sources and goals during speech planning, our study is the first to prove that Chinese exhibits a fragile goal bias at both the linguistic level and the nonlinguistic level.

5. Conclusion and future directions

The present study investigated the effects of audience design and goal bias in Chinese speakers' message generation of source-goal motion events. We verified that: (1) speakers considered listeners' knowledge state in their message generation of source and goal in source-goal motion event descriptions is a language-general phenomenon; (2) Chinese speakers show language-specific characteristics in their message generation of source and goal in motion events (i.e., fragile goal bias).

Some studies indicated that the profile of the addressee affects how the speaker models the addressee's knowledge state. For example, speakers were more likely to use redundant references with language learners (Tal et al., 2023). Adults and children tended to add information about instruments when communicating with an interactive listener (Grigoroglou & Papafragou, 2019). Future studies may explore to what extent the profile of the addressee affects speakers' message generation of source-goal motion events.

Previous studies have mainly investigated how crosslinguistic differences in motion event expressions influence speakers' conceptualization of manner and path (Ji & Hohenstein, 2017; Zhao & Hu, 2018), and just a few studies explored how the crosslinguistic difference in the expression of path elements, such as Trajectory and Goal affects speakers' message generation of path of motion (Flecken et al., 2014; Liao et al., 2020). Future research may explore whether satellite-framed language like English may differ from verb-framed languages like Spanish in speakers' message generation of paths of motion, including Source, Trajectory and Goal.

Data availability statement. The materials, data and analysis scripts are available through the Open Science Framework (https://osf.io/bywu5/?view_only=519cd8cea7da4de09eb2ac97f35c3598).

Funding statement. This work was supported by the MOE Project at Center for Linguistics and Applied Linguistics, Guangdong University of Foreign Studies (22JJD740021).

References

- Arnold, J. E. (2001). The effect of thematic roles on pronoun use and frequency of reference continuation. *Discourse Processes*, 31(2), 137–162. https://doi.org/10.1207/S15326950DP3102_02.
- Arnold, J. E. (2008). Reference production: Production-internal and addressee-oriented processes. *Language and Cognitive Processes*, 23(4), 495–527. <https://doi.org/10.1080/01690960801920099>.
- Barr, D. J., Levy, R., Scheepers, C., & Tily, H. J. (2013). Random effects structure for confirmatory hypothesis testing: Keep it maximal. *Journal of Memory and Language*, 68(3), 255–278. <https://doi.org/10.1016/j.jml.2012.11.001>.
- Bock, K. (1995). Chapter 6—Sentence production: From mind to mouth. In J. L. Miller & P. D. Eimas (Eds.), *Speech, language, and communication* (2nd ed., pp. 181–216). Academic Press. <https://doi.org/10.1016/B978-012497770-9.50008-X>
- Bock, K., & Levelt, W. (1994). Language production: Grammatical encoding. In M. A. Gernsbacher (Ed.), *Handbook of psycholinguistics* (pp. 945–984). Academic Press.

- Brennan, S. E., & Clark, H. H. (1996). Conceptual pacts and lexical choice in conversation. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 22(6), 1482–1493. <https://doi.org/10.1037/0278-7393.22.6.1482>.
- Brennan, S. E., Galati, A., & Kuhlen, A. K. (2010). Chapter 8—Two minds, one dialog: Coordinating speaking and understanding. In B. H. Ross (Ed.), *Psychology of learning and motivation* (Vol. 53, pp. 301–344). Academic Press. [https://doi.org/10.1016/S0079-7421\(10\)53008-1](https://doi.org/10.1016/S0079-7421(10)53008-1).
- Brown-Schmidt, S., & Konopka, A. (2008). Little houses and casas pequeñas: Message formulation and syntactic form in unscripted speech with speakers of English and Spanish. *Cognition*, 109, 274–280. <https://doi.org/10.1016/j.cognition.2008.07.011>.
- Brown-Schmidt, S., & Tanenhaus, M. K. (2006). Watching the eyes when talking about size: An investigation of message formulation and utterance planning. *Journal of Memory and Language*, 54(4), 592–609. <https://doi.org/10.1016/j.jml.2005.12.008>.
- Brown, P. M., & Dell, G. S. (1987). Adapting production to comprehension: The explicit mention of instruments. *Cognitive Psychology*, 19(4), 441–472. [https://doi.org/10.1016/0010-0285\(87\)90015-6](https://doi.org/10.1016/0010-0285(87)90015-6).
- Brown, V. A. (2021). An introduction to linear mixed-effects modeling in R. *Advances in Methods and Practices in Psychological Science*, 4(1), 2515245920960351. <https://doi.org/10.1177/2515245920960351>.
- Chen, L., & Guo, J. (2009). Motion events in Chinese novels: Evidence for an equipollently-framed language. *Journal of Pragmatics*, 41(9), 1749–1766. <https://doi.org/10.1016/j.pragma.2008.10.015>.
- Chen, Y., Trueswell, J., & Papafragou, A. (2024). Sources and goals in memory and language: Fragility and robustness in event representation. *Journal of Memory and Language*, 135, 104475. <https://doi.org/10.1016/j.jml.2023.104475>.
- Clark, H. H. (1996). *Using language*. Cambridge University Press.
- Clark, H. H., & Carlson, T. B. (1982). Context for comprehension. In J. Long & A. Baddeley (Eds.), *Attention and performance* (pp. 1–37). Academic Press.
- Do, M. L., Papafragou, A., & Trueswell, J. (2020). Cognitive and pragmatic factors in language production: Evidence from source-goal motion events. *Cognition*, 205, 104447. <https://doi.org/10.1016/j.cognition.2020.104447>.
- Do, M. L., Papafragou, A., & Trueswell, J. (2022). Encoding motion events during language production: Effects of audience design and conceptual salience. *Cognitive Science*, 46(1), e13077. <https://doi.org/10.1111/cogs.13077>.
- Flecken, M., Carroll, M., Weimar, K., & Von Stutterheim, C. (2015). Driving along the road or heading for the village? Conceptual differences underlying motion event encoding in French, German, and French-German L2 users. *The Modern Language Journal*, 99(S1), 100–122. <https://doi.org/10.1111/j.1540-4781.2015.12181.x>.
- Flecken, M., Von Stutterheim, C., & Carroll, M. (2014). Grammatical aspect influences motion event perception: Findings from a cross-linguistic non-verbal recognition task. *Language and Cognition*, 6(1), 45–78. <https://doi.org/10.1017/langcog.2013.2>.
- Grice, H. P. (1975). Logic and conversation. In P. Cole & J. Morgan (Eds.), *Syntax and semantics 3: Speech acts* (pp. 41–58). Academic Press.
- Grigoroglou, M., & Papafragou, A. (2019). Children's (and adults') production adjustments to generic and particular listener needs. *Cognitive Science*, 43(10), e12790. <https://doi.org/10.1111/cogs.12790>.
- Horton, W. S., & Gerrig, R. J. (2002). Speakers' experiences and audience design: Knowing when and knowing how to adjust utterances to addressees. *Journal of Memory and Language*, 47(4), 589–606. [https://doi.org/10.1016/S0749-596X\(02\)00019-0](https://doi.org/10.1016/S0749-596X(02)00019-0).
- Horton, W. S., & Gerrig, R. J. (2016). Revisiting the memory-based processing approach to common ground. *Topics in Cognitive Science*, 8(4), 780–795. <https://doi.org/10.1111/tops.12216>.
- Horton, W. S., & Keysar, B. (1996). When do speakers take into account common ground? *Cognition*, 59(1), 91–117. [https://doi.org/10.1016/0010-0277\(96\)81418-1](https://doi.org/10.1016/0010-0277(96)81418-1).
- Ji, Y., & Hohenstein, J. (2017). Conceptualising voluntary motion events beyond language use: A comparison of English and Chinese speakers' similarity judgments. *Lingua*, 195, 57–71. <https://doi.org/10.1016/j.lingua.2017.05.007>.
- Johanson, M., Selimis, S., & Papafragou, A. (2019). The source-goal asymmetry in spatial language: Language-general vs. language-specific aspects. *Language, Cognition and Neuroscience*, 34(7), 826–840. <https://doi.org/10.1080/23273798.2019.1584323>.

- Konopka, A. E., & Brown-Schmidt, S. (2014). Message encoding. In V. Ferreira, M. Goldrick, & M. Miozzo (Eds.), *The oxford handbook of language production* (pp. 3–20). Oxford University Press.
- Kuhlen, A. K., & Brennan, S. E. (2010). Anticipating distracted addressees: How speakers' expectations and addressees' feedback influence storytelling. *Discourse Processes*, 47(7), 567–587. <https://doi.org/10.1080/01638530903441339>.
- Lakusta, L., & Carey, S. (2015). Twelve-month-old infants' encoding of goal and source paths in agentive and non-agentive motion events. *Language Learning and Development*, 11(2), 152–157. <https://doi.org/10.1080/15475441.2014.896168>.
- Lakusta, L., & DiFabrizio, S. (2017). And, the winner is...a visual preference for endpoints over starting points in infants' motion event representations. *Infancy*, 22(3), 323–343. <https://doi.org/10.1111/infa.12153>.
- Lakusta, L., & Landau, B. (2005). Starting at the end: The importance of goals in spatial language. *Cognition*, 96(1), 1–33. <https://doi.org/10.1016/j.cognition.2004.03.009>.
- Lakusta, L., & Landau, B. (2012). Language and memory for motion events: Origins of the asymmetry between source and goal paths. *Cognitive Science*, 36(3), 517–544. <https://doi.org/10.1111/j.1551-6709.2011.01220.x>.
- Lakusta, L., Muentener, P., Petrillo, L., Mullanaphy, N., & Muniz, L. (2016). Does making something move matter? Representations of goals and sources in motion events with causal sources. *Cognitive Science*, 41(3), 814–826. <https://doi.org/10.1111/cogs.12376>.
- Lakusta, L., Spinelli, D., & Garcia, K. (2017). The relationship between pre-verbal event representations and semantic structures: The case of goal and source paths. *Cognition*, 164, 174–187. <https://doi.org/10.1016/j.cognition.2017.04.003>.
- Lakusta, L., Wagner, L., O'Hearn, K., & Landau, B. (2007). Conceptual foundations of spatial language: Evidence for a goal bias in infants. *Language Learning and Development*, 3(3), 179–197. <https://doi.org/10.1080/15475440701360168>.
- Levelt, W. J. M. (1989). *Speaking: From intention to articulation*. The MIT Press.
- Liao, Y., Flecken, M., Dijkstra, K., & Zwaan, R. A. (2020). Going places in Dutch and mandarin Chinese: Conceptualising the path of motion cross-linguistically. *Language, Cognition and Neuroscience*, 35(4), 498–520. <https://doi.org/10.1080/23273798.2019.1676455>.
- Liu, X., & Wen, Q. (2023). Judgment method of conceptualization transfer: An eye-tracking study of motion events based on cognitive contrastive analysis. *Foreign Language Teaching and Research*, 55(2), 212–319. <https://doi.org/10.19923/j.cnki.fltr.2023.02.006>.
- Lockridge, C. B., & Brennan, S. E. (2002). Addressees' needs influence speakers' early syntactic choices. *Psychonomic Bulletin & Review*, 9(3), 550–557. <https://doi.org/10.3758/BF03196312>.
- Pace, A., Levine, D. F., Golinkoff, R. M., Carver, L. J., & Hirsh-Pasek, K. (2020). Keeping the end in mind: Preliminary brain and behavioral evidence for broad attention to endpoints in pre-linguistic infants. *Infant Behavior and Development*, 58, 101425. <https://doi.org/10.1016/j.infbeh.2020.101425>.
- Papafragou, A. (2010). Source-goal asymmetries in motion representation: Implications for language production and comprehension. *Cognitive Science*, 34(6), 1064–1092. <https://doi.org/10.1111/j.1551-6709.2010.01107.x>.
- Papafragou, A., & Grigoroglou, M. (2019). The role of conceptualization during language production: Evidence from event encoding. *Language, Cognition and Neuroscience*, 34(9), 1117–1128. <https://doi.org/10.1080/23273798.2019.1589540>.
- R Core Team. (2018). *R: A language and environment for statistical computing*. R Foundation for Statistical Computing. <https://www.R-project.org/>
- Regier, T. (1996). *The human semantic potential: Spatial language and constrained connectionism*. MIT Press.
- Regier, T., & Zheng, M. (2007). Attention to endpoints: A cross-linguistic constraint on spatial meaning. *Cognitive Science*, 31(4), 705–719. <https://doi.org/10.1080/15326900701399954>.
- Slobin, D. I. (1996). From “thought and language” to “thinking for speaking”. In J. J. Gumperz & S. C. Levinson (Eds.), *Rethinking linguistic relativity* (pp. 70–96). Cambridge University Press.
- Stefanowitsch, A., & Rohde, A. (2004). The goal bias in the encoding of motion events. In G. Radden & K.-U. Panther (Eds.), *Studies in linguistic motivation* (pp. 249–267). Mouton de Gruyter.
- Tal, S., Grossman, E., Rohde, H., & Arnon, I. (2023). Speakers use more redundant references with language learners: Evidence for communicatively-efficient referential choice. *Journal of Memory and Language*, 128, 104378. <https://doi.org/10.1016/j.jml.2022.104378>.

- Talmy, L. (1985). Lexicalization patterns: Semantic structure in lexical forms. In T. Shopen (Ed.), *Grammatical categories and the lexicon. Language typology and syntactic description* (Vol. 3). Cambridge University Press.
- Van Der Wege, M. M. (2009). Lexical entrainment and lexical differentiation in reference phrase choice. *Journal of Memory and Language*, 60(4), 448–463. <https://doi.org/10.1016/j.jml.2008.12.003>.
- Vanlangendonck, F., Willems, R. M., Menenti, L., & Hagoort, P. (2016). An early influence of common ground during speech planning. *Language, Cognition and Neuroscience*, 31(6), 741–750. <https://doi.org/10.1080/23273798.2016.1148747>.
- von Stutterheim, C., & Nüse, R. (2003). Processes of conceptualization in language production: Language-specific perspectives and event construal. *Language*, 41(5), 851–881. <https://doi.org/10.1515/ling.2003.028>.
- Zhang, S. (2021). Directions and domains of conceptual transfer: Evidence from motion events. *Foreign Language Teaching and Research*, 53(3), 400–412. <https://doi.org/10.19923/j.cnki.fltr.2021.03.007>.
- Zhao, C., & Hu, B. (2018). The role of event structure in language production: Evidence from structural priming in Chinese motion event descriptions. *Lingua*, 208, 61–81. <https://doi.org/10.1016/j.lingua.2018.03.006>.
- Zhao, C., & Li, J. (2022). Effects of cross-language differences on memory of motion events in English and Chinese. *Modern Foreign Languages*, 45(1), 102–133. <https://link.cnki.net/urlid/44.1165.H.20210919.1114.018>
- Zheng, M., & Goldin-Meadow, S. (2002). Thought before language: How deaf and hearing children express motion events across cultures. *Cognition*, 85(2), 145–175. [https://doi.org/10.1016/S0010-0277\(02\)00105-1](https://doi.org/10.1016/S0010-0277(02)00105-1).

Cite this article: Zhao, C., Xu, R., & Sun, T. (2025). The role of audience design and goal bias in message generation: Evidence from Chinese source-goal motion events, *Language and Cognition*, 17, e65, 1–25. <https://doi.org/10.1017/langcog.2025.10024>