

From Multimodal Perception to Strategic Reasoning: A Survey on AI-Generated Game Commentary

Qirui Zheng^{1,*}, Xingbo Wang^{1,*}, Keyuan Cheng^{1,*}, Yunlong Lu¹, Muhammad Asif Ali², Lingfeng Li¹, Yongyi Wang¹ and Wenxin Li^{1,✉}

¹Peking University

²King Abdullah University of Science and Technology
zzzqr@stu.pku.edu.cn, jacksymbol@stu.pku.edu.cn, lwx@pku.edu.cn

Abstract

The advent of artificial intelligence has propelled AI-Generated Game Commentary (AI-GGC) into a rapidly expanding research area, offering advantages such as scalable availability and personalized narration. However, existing studies remain fragmented, and a systematic survey that unifies prior efforts is still lacking. To bridge this gap, our survey introduces a unified framework that systematically organizes the AI-GGC landscape. We present a novel taxonomy focused on three core commentator capabilities: Live Observation, Strategic Analysis, and Historical Recall, and further categorize commentary into three corresponding types: Descriptive Commentary, Analytical Commentary, and Background Commentary. Building on this structure, we provide an in-depth review of methods, datasets, and evaluation metrics, analyzing their strengths and limitations. Finally, we highlight key challenges and point out promising directions for future research in AI-GGC.

1 Introduction

Game commentary is the real-time narration and analysis of gameplay, which enhances the viewing experience by elucidating strategies and highlighting pivotal moments [Jhamtani *et al.*, 2018]. Game commentary often appear in three major domains: board games (Chess, Go), sports (basketball, soccer), and esports (League of Legends (LoL), Dota 2). Its importance is evidenced by the massive viewership of major tournaments, such as the World Cup Final, which attract billions of concurrent viewers and rely heavily on commentary to sustain engagement and immersion.

The rapid advancement of AI has made *AI-Generated Game Commentary* (AI-GGC) feasible. AI commentators present several advantages over human commentators. Primarily, they ensure unlimited availability, covering matches that cannot be covered by human commentators due to limited broadcast resources [Siu *et al.*, 2023]. Furthermore, they offer unparalleled personalization [Andrews *et al.*, 2024a]. Through tailored training, AI commentators can adapt to diverse viewer preferences, including the choice of language, voice, favored teams, and the level of analytical detail.

Although considerable research has been conducted, the field still lacks a systematic survey that consolidates these fragmented endeavors. This fragmentation stems primarily from two fundamental challenges inherent to the domain:

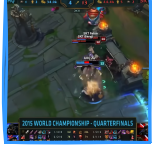
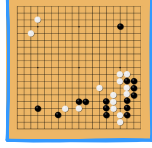
(1) Strategic depth parsing. In board games such as Chess and Go, while RL-based AI systems have achieved superhuman performance [Silver *et al.*, 2016], pure Large Language Models (LLMs) often struggle to provide deep, principled analysis of the reasoning behind each move. They can describe the moves made but frequently fail to articulate the underlying strategic reasoning and long-term planning behind them—reflecting their limited capacity for complex logical reasoning and long-horizon strategic inference.

(2) Multi-modal heterogeneity. Different game genres essentially “speak different languages”. Board games rely on simple, symbolic formats; sports games are captured as fast-paced, visually rich videos; while e-sports generate high-frequency, complex metadata streams. The entirely different techniques needed to process each format thus segregate research into isolated camps, preventing a unified methodology.

To address this gap, Section 2 introduces a systematic survey scheme for analyzing AI commentators, guided by two key questions: “*What makes a proficient commentator?*” and “*What defines effective commentary?*” We clarify the three core capabilities required of AI commentators and propose a taxonomy that classifies commentary into three types by content and purpose. Building on this foundation, Section 3 reviews technical methodologies for these core capabilities. Section 4 surveys datasets by game genre and data practices, while Section 5 analyzes evaluation metrics. Figure 2 presents our integrated taxonomy of methods, datasets, and metrics. Finally, Section 6 discusses open challenges and future directions, and Section 7 concludes the paper. To summarize, this survey makes following key contributions:

1. We propose a unified survey scheme that organizes AI-GGC along game genres, core commentator capabilities, and commentary types, providing a principled structure for analyzing existing methods and their trade-offs.
2. Building on this scheme, we present the first comprehensive survey of AI-GGC spanning methods, datasets, and evaluation, and introduce a systematic taxonomy that synthesizes fragmented literature into a coherent framework.
3. We critically examine the current AI-GGC landscape,

(a) Game Genres



(b) Core Capabilities of AI Commentators

Live Observation



Strategic Analysis



Historical Recall

	Team	Games	Wins	Losses	Win %	K/D	CS	DM	DM%	DM%	DM%	DM%
1	TLX	12	12	0	100%	17.6	57,962	0.85	2.2	9.1	89.2%	89.2%
2	CFO	12	10	2	75%	12.5	49,973	0.69	2.1	7.8	46.2%	76.9%
3	OCO	12	8	4	66.7%	16.2	54,481	0.82	2.2	6.8	46.2%	76.9%
4	FLK	12	0	12	0%	12.8	56,881	0.84	2.4	6.1	62.5%	66.5%

(c) Commentary Types

Descriptive Commentary

Messi receives the ball in the center circle. He turns and plays a quick pass out to the winger on the right touchline.

Analytic Commentary

A classic trade, but White gets the bishop pair. That's a long-term strategic advantage in this open position.

Background Commentary

A scene we've seen since 2013. Just when it looks darkest, Faker steps up to rewrite the fight, a timeless classic from the GOAT.

Figure 1: Overview of the proposed AI-GGC survey scheme, systematically summarizing the field along three key dimensions: (a) Game genres; (b) Foundational capabilities of AI commentators; (c) Commentary types.

highlighting challenges in capability coverage and maturity, system-Level coordination, and evaluation.

2 AI-GGC Survey Scheme

In this section, we propose a systematic survey scheme, illustrated in Figure 1, which organizes and summarizes the field of AI-GGC along three principal dimensions: game genre, foundational AI capabilities, and commentary types.

2.1 Scope and Game Genres

Our survey scheme is anchored in three representative and distinct game genres: (1) **Board Games**. Covers games with perfect information (Chess) or imperfect information (Mahjong). Commentary focuses on reasoning, move quality evaluation, and strategic analysis. (2) **Esports**. Includes fast-paced digital titles like Dota and LoL. Commentary analyzes real-time decision-making and both individual and team-based strategies. (3) **Sports**. Physical, dynamic games like soccer and basketball. Commentary emphasizes real-time movement, key events, and tactical coaching decisions.

In Section 4, we categorize and analyze existing AI-GGC datasets according to these game genres.

2.2 Core Capabilities of AI Commentators

Drawing inspiration from the expertise of human commentators across various game genres, we abstract the core cognitive modules that underpin high-quality AI-GGC systems.

(1) **Live Observation (LO)** is responsible for the accurate perception and interpretation of real-time events and fine-grained details within the game environment. It serves as the primary data source for generating descriptive commentary, such as identifying the placement of pieces in chess, or describing shots and passes in soccer, kills and assists in LoL.

(2) **Strategic Analysis (SA)** encompasses the analytical abilities required to comprehend the current game state, interpret strategic intent, evaluate decisions, and anticipate future

developments. It is particularly vital in board games, where understanding the intention and implications of each move is essential for insightful commentary.

(3) **Historical Recall (HR)** involves leveraging extensive game knowledge, including player histories, classic match moments, and community-specific terminology. It enables the commentator to provide meaningful context, such as recognizing signature moves or recalling iconic past plays, thereby enhancing the informativeness and engagement.

A systematic discussion of how existing methods implement these foundational capabilities is provided in Section 3.

2.3 Taxonomy of Commentary Types

According to the three cognitive modules, we classify commentary types, which provides a new evaluation dimension.

(1) **Descriptive Commentary**. This is the most fundamental type, directly derived from the LO capability. It objectively addresses the question, “*What is happening?*” by narrating immediate, observable events in games.

(2) **Analytical Commentary**. Derived from SA, this commentary transcends basic description by evaluating the quality of decisions or actions (“*How good?*”), interpreting player intent (“*Why?*”), and predicting strategic outcomes within a broader tactical context (“*What’s next?*”).

(3) **Background Commentary**. Enabled by the HR capability, this commentary type enriches the live narrative by integrating broader contextual information, such as detailed team and player backgrounds, game history, and cultural references. It addresses the question “*What is the significance of this?*” by referencing past performances, historical match situations, team history, and community knowledge.

3 Methods

This section reviews technical pathways towards core capabilities of AI commentators introduced in Section 2. We also

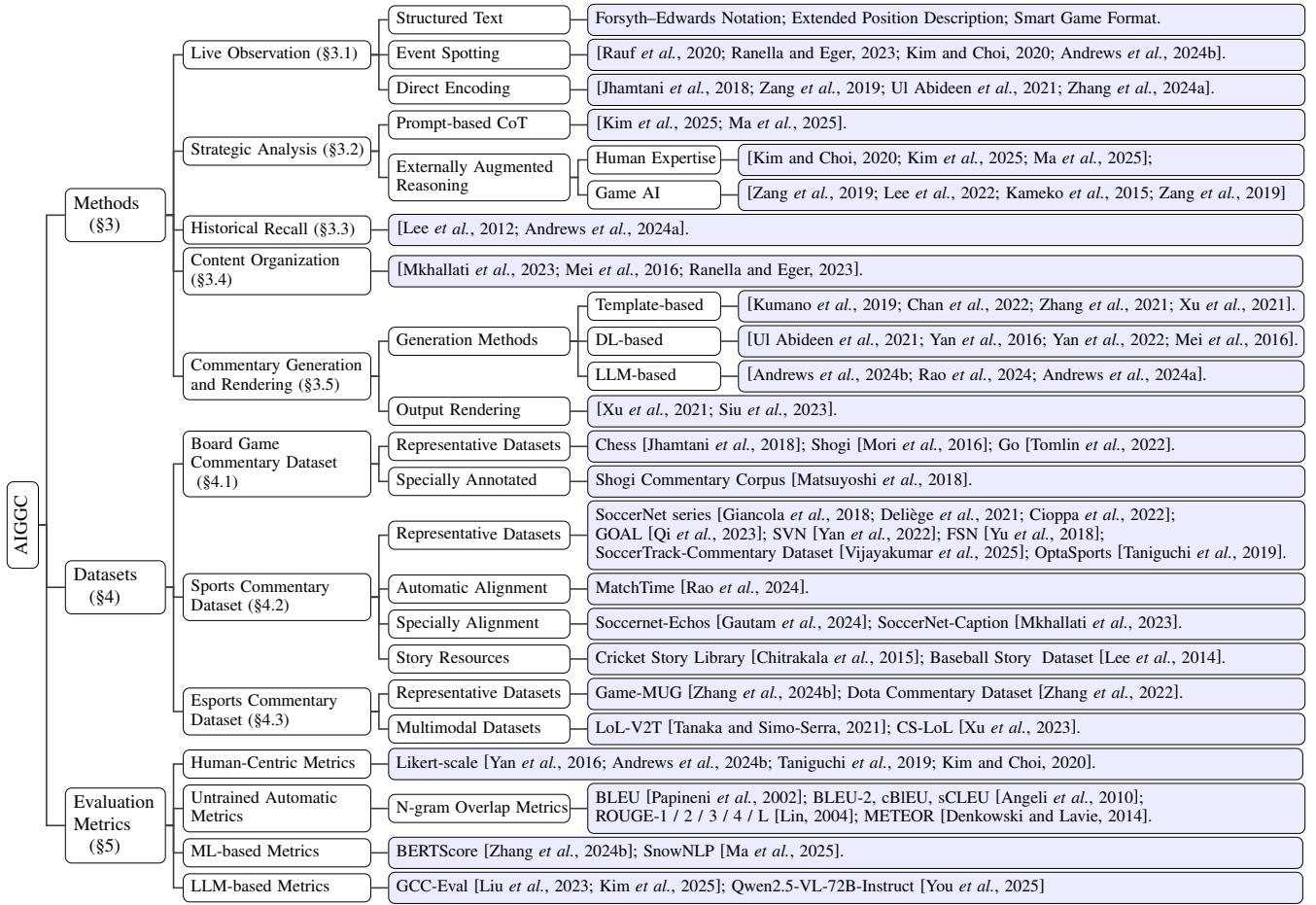


Figure 2: Taxonomy of Methods, Datasets and Metrics in AI-GGC

expand our discussions to architectural designs that wire key function modules to produce end-to-end game commentaries.

3.1 Live Observation

LO requires AI-GGC systems to process complex, multi-modal game data under time constraints. Board games typically rely on direct visual observation of game board or structured textual representation that encode game state into compact strings (e.g. *Forsyth-Edwards Notation* for chess), while sports and e-sports rely mostly on videos. However, explaining “what is going on” in fast-paced games involves more than multimodal perception. A good commentary should be able to identify focal points in chaotic gameplays and summarize them using professional terminology, balancing detail and conciseness under time and comprehension constraints

Early technical trials attempting to address this challenge primarily adopt **Event Spotting (ES)**, which reduces the complexity of raw game video processing by converting visual inputs into textual labels. Specifically, predefined event labels are assigned to corresponding timestamps in game videos and subsequently used as inputs to the commentary generation model, instead of raw video streams. Early works [Giancola et al., 2018; Deliège et al., 2021] rely on human-annotated domain-specific datasets to train models for

ES which may suffer from generalization issues. Detection-based ES reduces annotation labor by using object-detection models to identify actions and entities in videos. Jung [2025] detect players and the ball using YOLOv8 and perform ES by combining spatial-temporal trajectories, action recognition, and rule-based reasoning to identify key basketball events. Similar pipelines have also been explored using various object detection backbones [Ranella and Eger, 2023].

While effective in reducing perceptual complexity, this representation inevitably introduces information loss by discarding fine-grained details beyond the detected events.

Direct Encoding approaches map multimodal inputs (e.g. video streams and metadata) into dense representations in order to preserve information completeness. The advancing multimodal large language models has made it possible. Initial efforts focused on encoding structured representations for board games, directly encoding chessboard states into vectors [Jhamtani et al., 2018; Zang et al., 2019], until recent advance in models’ multimodal understanding capabilities enable the extension of this approach to visual inputs from sports. Jhamtani [2024a] combined video transformers for visual input and BERT for textual descriptions, fusing these features with a transformer encoder guided by a soft-prompt mechanism for better multimodal perception.

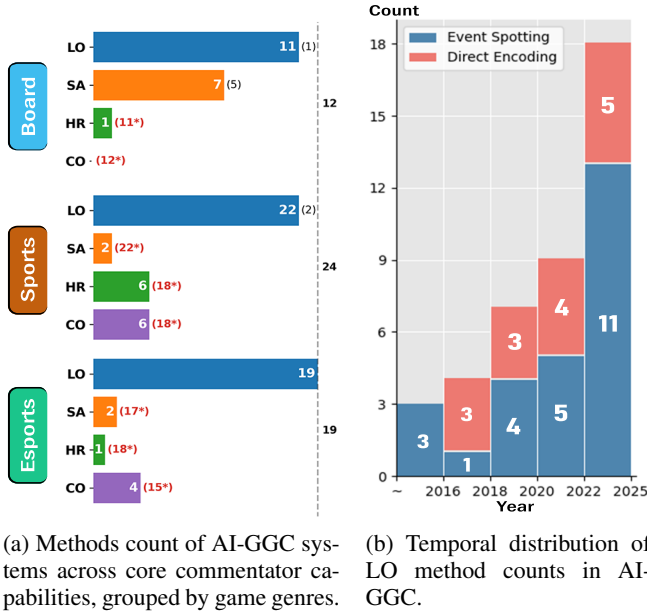


Figure 3: Methodological patterns in AI-GGC research across capabilities and time. All statistics are reported by paper count.

Despite recent progress, several direct encoding approaches still rely on relatively early visual encoding backbones [Ghosh *et al.*, 2023], which may limit their capabilities for modeling complex visual information. More fundamentally, recent studies show that even state-of-the-art LLMs struggle with visual perception in low-dynamic board games [Wang *et al.*, 2025], raising concerns about their scalability to fast-paced and highly dynamic sports or esports.

Methodologically, LO paradigms in AI-GGC have evolved in response to changing modeling constraints and opportunities, as reflected in the temporal distribution of methods in Figure 3b. Early work almost exclusively adopted ES to simplify video modeling under limited capacity, inevitably introducing information loss. With advances in DL, direct encoding emerged as a parallel paradigm between 2016 and 2022, aiming to preserve richer visual information. Since 2023, the rise of LLMs has renewed interest in event-based representations, as discrete events can be readily transformed into textual or structured prompts without resolving visual perception, making ES attractive from an engineering perspective. Importantly, these trends do not imply that direct encoding is a dead end. Rather, its practicality remains contingent on advances in multimodal understanding, suggesting that it may regain relevance as model capabilities continue to evolve.

3.2 Strategic Analysis

A key challenge for AI-GGC is enabling deep SA, a capability that places substantial demands on the logical reasoning ability. Beyond LO, a good commentator is expected to address “How good”, explain “Why”, and anticipate “What’s next” within the same logic. Depending on whether SA relies solely on the model itself, we categorize existing AI-GGC methods for SA into *Prompt-based Chain-of-Thought (CoT)* and *Externally Augmented Reasoning* [Li *et al.*, 2025b].

Prompt-based CoT refers to eliciting step-by-step SA purely through instructions in the prompt, guiding the model to carry out deeper reasoning without introducing external modules, search procedures, or additional training. In AI-GGC, prompting is not only used to steer how the model analyzes the current situation, but also to inject details of game rules and gameplay common sense that LLMs may be unfamiliar with (e.g., in Guandan, three-with-a-pair is allowed whereas three-with-a-single is not). Prior work [Kim *et al.*, 2025; Ma *et al.*, 2025] has shown that such prompt-based guidance can better elicit SA ability of LLMs and improve the quality of generated commentary.

Despite its simplicity and effectiveness, this method mainly elicits SA already latent in the model, without extending its intrinsic reasoning capacity. As a result, its ability to support deeper strategic inference and forward-looking anticipation remains bounded by the model itself.

Externally Augmented Reasoning refers to enhancing SA by introducing auxiliary components beyond the language model itself, such as advanced algorithms or expert knowledge, to compensate for limited reasoning capacity of models. Such augmentation typically takes the form of external guidance, which can be broadly categorized into guidance derived from *Human Expertise* and guidance provided by *Game AIs*.

Incorporating strategic principles from **Human Expertise** allows LLMs to perform more insightful analysis, analogous to expert commentators explaining key moves and tactical intentions. The main challenge lies in embedding such principles effectively into the model’s reasoning process. Kim [2020] addressed this by manually mapping baseball events to strategic concepts and provided both as model inputs, enabling interpretation within an explicit strategic context. In chess, Kim [2025] extracted key strategic concepts (e.g., king safety) using a game engine and compared their representations before and after each move to identify salient strategic changes for commentary generation. For imperfect-information games such as Guandan, Ma [2025] proposed a two-level prompting scheme grounded in Theory of Mind (ToM) [Frith and Frith, 2005], where one prompt guided the model to analyze its own strategy while the other encouraged inference over opponents’ possible hands and intentions.

Augmenting strategy analysis with powerful **Game AIs** is another natural idea, mirroring common practice among human commentators, particularly in some board games where game AIs are dominant. However, this approach introduces additional challenges. Despite their strength, RL-based game AIs are often black boxes, limiting interpretability and making it difficult to directly guide logical reasoning. Consequently, existing approaches focus on transforming unexplainable outputs into more explainable signals. An early idea is to achieve this through controllable search or rule-based mechanisms. Kameko [2015] employed rule-based scoring functions and limited-depth search in Shogi to generate candidate game trees and derive commentary from simulated outcomes. Another idea is to translate engine outputs into interpretable linguistic signals. Lee [2022] quantified “move quality” by measuring win rate difference between the played move and the AI-recommended best move, and this difference was mapped to discrete natural-language categories. Zhang [2019] jointly

trained a chess engine and a language model in a multi-task framework, allowing the system to simultaneously learn game understanding and commentary generation.

Although externally augmented reasoning substantially improves the logical coherence and correctness of analytical commentary, it also faces inherent limitations. First, both sources of external guidance exhibit inherent weaknesses. Human expertise provides only partial coverage of the strategic space, generalizes poorly to complex or unseen situations, and may become outdated over time. Game AI-based augmentation, while powerful, is constrained by the black-box nature of modern engines; simple mappings from engine outputs to linguistic labels are often insufficient for rich strategic explanations. More critically, in highly complex sports and esports domains, dominant game AIs are still largely unavailable, severely limiting the applicability of AI-driven augmentation. Second, external augmentation introduces capabilities absent from the model itself, but these remain externally provided rather than internalized as intrinsic reasoning ability.

Beyond method-specific limitations, existing AI-GGC approaches share a fundamental weakness in SA: they mainly focus on short-horizon quality assessment (“How good?”), without constructing coherent multi-step strategic narratives. As a result, reasoning about “Why?” is often incomplete, leading to unreliable prediction of “what’s next” based on such logic. And as also discussed above, current approaches rarely internalize strategic reasoning into the model itself, instead mainly relying on externally provided signals.

3.3 Historical Recall

Historical recall enriches background commentary by integrating external knowledge. Previous works have approached this primarily as an information retrieval task. Lee [2012] ranked stories from a database based on their similarity to the current game state using AdaRank [Xu and Li, 2007], then filtered candidates with a regression model before inserting into commentary. Andrews [2024a] encoded player information with SBERT to retrieve relevant player or team data.

3.4 Content Organization

While techniques in former sections have enabled rich contents and details based on raw game states and data, real game commentary is usually restricted by limited time and audiences’ capabilities to absorb the game. Thus AI-GGC faces an inherent challenge: transforming unstructured segments into logically structured, concise narratives through strategic selection and prioritization to maintain viewer comprehension without cognitive overload. Mkhallati [2023] proposed a spotting head that predicts when commentary should be generated for each time window, using non-maximum suppression to eliminate redundancy. Mei [2016] introduced a coarse-to-fine aligner that filters out irrelevant records and reweights alignment scores to focus on salient events. Likewise, Ranella [2023] used an event-priority and delayed-confirmation mechanism to filter key events in real time.

3.5 Commentary Generation and Rendering

Commentary generation, the final stage of AI-GGC, transforms processed information into natural language. These

methods fall into three main categories: template-based, DL-based, and LLM-based. Template-based approaches are efficient and reliable [Kumano *et al.*, 2019; Chan *et al.*, 2022; Zhang *et al.*, 2021], but often yield repetitive outputs. DL-based methods, typically using LSTM decoders [Yan *et al.*, 2016; Yan *et al.*, 2022; Mei *et al.*, 2016], provide more diverse and context-aware commentary. LLMs further improve fluency and strategic reasoning, with prompt-based techniques explored in recent work [Andrews *et al.*, 2024b; Andrews *et al.*, 2024a; Kim *et al.*, 2025], though they can be slower, more costly, and prone to hallucinations.

Multimodal rendering is an emerging area, with most non-textual commentaries generated by converting text to speech via TTS APIs. For example, Xu [2021] generated both text and audio parameters for expressive speech, while Siu [2023] combined TTS and StyleGAN to create a virtual commentator with realistic facial expressions and lip-syncing.

3.6 A Systematic Gap in Capability Coverage

Figure 3a reveals a pronounced structural gap in current AI-GGC research: most systems focus on LO and fail to cover the full set of core commentator capabilities, with comparatively limited attention to SA, HR, and CO. This imbalance manifests differently across game genres due to their distinct characteristics. For board games, longer turn intervals and high demands for logical reasoning lead to greater research attention on SA, while reducing the relative need for CO. Powerful game AIs provide strong support for this ability, reinforcing this focus. In contrast, the limited attention to HR is difficult to justify. For sports and esports, the dominance of LO is largely driven by the intrinsic difficulty of multimodal perception in fast-paced and visually complex scenes, leaving other capabilities substantially under-explored despite their importance for effective commentary.

4 Datasets and Benchmarks

This section reviews dataset according to game genres, along with relevant data processing techniques.

4.1 Board Game Commentary Datasets

Commentary datasets for board games have been developed for chess, Go, Shogi (Japanese chess), GuanDan (Chinese card game). The turn-based nature of board games ensures natural alignment between game states and commentary, and the slower pace offers sufficient time for deep reasoning.

Most datasets consist of aligned pairs of game states and commentaries. Jhamtani [2018] introduced a large chess dataset with 298K move-by-move annotations, establishing a benchmark for AI-GGC in chess. Additional datasets exist for Shogi [Kameko *et al.*, 2015] and Go [Tomlin *et al.*, 2022].

Some datasets are annotated beyond commentary text for better commenting. A Shogi commentary corpus [Matsuyoshi *et al.*, 2018] annotates modality expressions and event factuality, enabling fine-grained modeling of uncertainty and future-oriented reasoning in strategic commentary.

4.2 Sports Commentary Datasets

As mentioned before, in sport games, ES is widely used to simplify perception of complex multimodal input. ES-based

datasets cover soccer [Qi *et al.*, 2023], basketball [Yan *et al.*, 2022], baseball [Kim and Choi, 2020], and tennis [Yan *et al.*, 2016]. SoccerNet [Giancola *et al.*, 2018] is a prominent example, providing soccer videos with event labels and over 500K commentaries. It expanded to SoccerNet-v2 [Deliège *et al.*, 2021] with 17 event categories and SoccerNet-v3 [Cioppa *et al.*, 2022] with spatial annotations. To compensate for the information loss caused by ES, some datasets also incorporate additional modalities such as player tracking [Vijayakumar *et al.*, 2025] and play data [Taniguchi *et al.*, 2019], enabling richer perception and reasoning.

Manual alignment of game states and commentary remains common but is labor-intensive. To automate alignment, Rao [2024] developed a pipeline combining LLaMA-3 and contrastive learning, trained on a small set of aligned data.

Alternative alignment strategies have been introduced for specific tasks, including timestamp-based sparse alignment to support long-video commenting [Mkhallati *et al.*, 2023] and benchmarks where commentaries outnumber events, reflecting real-world settings [Zhang and Eickhoff, 2021].

Beyond commentary datasets, some auxiliary resources can be used to enrich background commentary by providing anecdotes, such as story libraries for cricket [Chitrakala *et al.*, 2015] and baseball [Lee *et al.*, 2012; Lee *et al.*, 2014].

4.3 Esports Commentary Datasets

Esports datasets, like those for sports, often use ES to manage complex video inputs. A unique advantage in esports is the availability of official APIs, which provide precise access to rich multimodal signals such as player statistics, game state metadata, and audio streams [Ringer *et al.*, 2019], placing higher demands on multimodal perception.

Esports commentary datasets cover games such as LoL [Wang and Yoshinaga, 2021], Dota 2 [Zhang *et al.*, 2022], and RoboCup [Liang *et al.*, 2009], and often include rich multimodal data including audio features [Zhang *et al.*, 2024b], sensor data [Chen *et al.*, 2010], or multiple view-points [Ishigaki *et al.*, 2023]. Ringer [2019] introduced affect labels capturing both game context and emotional dimensions, enabling affect-aware commentary generation.

5 Evaluation Metrics

Evaluating AI-GGC systems is inherently challenging for three reasons: (1) AI-GGC is a multi-faceted task involving various core commentator capabilities, with no clear way to aggregate them into a single score. (2) It is an open-ended NLG task with no gold reference, as different commentators may produce diverse yet equally valid commentaries for the same game state. (3) Analytical commentary often requires long chains of strategic reasoning that are difficult to formalize, making correctness hard to assess with existing metrics. Current AI-GGC work relies on evaluation methods that only partially address these challenges and have substantial limitations, which can be broadly grouped into four categories:

Human-Centric Metrics. Human-centric metrics remain the gold standard for assessing commentary quality, typically using questionnaires like Likert-scale. This metric of course can ignore the three challenges due to its directly reflecting

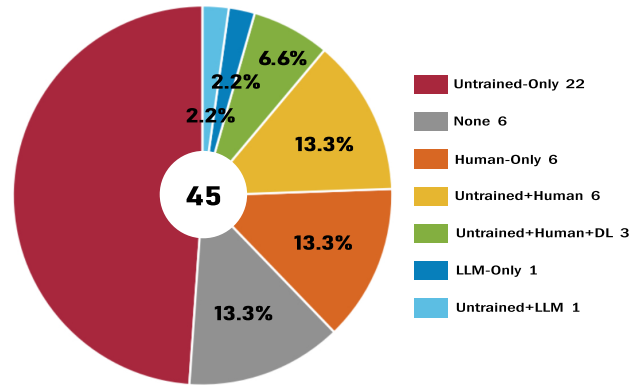


Figure 4: Distribution of evaluation strategies used in AI-GGC.

human preference, but they are expensive, require domain expertise, and highly subject to individual variability.

Untrained Automatic Metrics. Untrained automatic metrics, especially n-gram overlap measures, remain the most widely used tools for AI-GGC evaluation. As shown in Figure 3a, there are nearly half of AI-GGC works using it as the only evaluation metric. Metrics such as BLEU [Papineni *et al.*, 2002], ROUGE [Lin, 2004], and METEOR [Denkowski and Lavie, 2014] quantify lexical similarity to reference texts and are popular due to their low cost and reproducibility. However, n-gram overlap metrics are inherently misaligned with the task and correlate weakly with human judgments because AI-GGC lacks definitive ground-truth references [Reiter and Belz, 2009], making them ineffective for addressing any of the three challenges discussed above.

ML-based Metrics. ML-based metrics instead use learned models to approximate semantic similarity or other properties of commentary. Zhang [2024b] use BERTScore to evaluate esports commentary. However, it is still a reference-based measure of similarity to human texts and thus shares the same fundamental limitation as n-gram metrics in AI-GGC. Worth mentioning, Ma [2025] apply SnowNLP, a Naïve Bayes classifier, to estimate sentiment polarity of commentary. Although it is a very coarse signal and should not be treated as a primary metric, it has limited value as an auxiliary indicator of whether the overall emotional tone is plausible.

LLM-based Metrics move beyond simple similarity measures by providing reference-free, multi-dimensional assessments that better align with human preferences through prompting or post-training. For instance, Kim [2025] introduced GCC-Eval, which integrates expert chess knowledge with G-Eval [Liu *et al.*, 2023] to evaluate chess commentary along four dimensions: relevance, completeness, clarity, and fluency. However, it applies the same dimension weights to all commentary types and avoids assigning reasoning correctness to the LLM, relying instead on human annotations. Although still an emerging approach, LLM-as-a-judge currently appears to be the most promising direction for evaluating AI-GGC. However, existing work also shows that such judges remain unstable when verifying long chains of reasoning and exhibit systematic biases, such as a tendency to prefer longer [Ye *et al.*, 2024] and position bias [Shi *et al.*, 2024], which limits their current reliability for AI-GGC.

Overall, existing AI-GGC evaluation remains only partially aligned with the task: untrained and ML-based metrics struggle with open-endedness and multidimensionality, while LLM-based judges are still unreliable for assessing long-chain reasoning. To address these limitations, we argue that the taxonomy of descriptive, analytical, and background commentary offers a natural basis for type-specific evaluation, as discussed in Section 6.

6 Challenges and Future Directions

AI-GGC is rapidly emerging as a prominent application area for LLMs, where technical innovation meets growing real-world demand. While recent research and practical deployments highlight the field’s promise, they also expose fundamental challenges and point to key directions for future work.

Incomplete Coverage of Core Commentator Capabilities. A common limitation of existing AI-GGC systems is their incomplete coverage of core AI commentator capabilities, as shown in Figure 3a and discussed in Section 3.5. *Future AI-GGC research should therefore place balanced emphasis on all core commentator capabilities, treating them as equally important components in building more informative, interpretable, and complete commentary systems.*

Shallow Realization of LO. LO requires AI commentators to accurately capture and summarize ongoing games to guide audience attention and comprehension. Although a considerable amount of work has already been devoted to processing raw game data and extracting fundamental game events, these descriptions are still far from clear, specific, and easily understandable game descriptions. While pattern-recognition-based processing relies on models trained on annotated data with predefined event labels [Jung *et al.*, 2025], this class of methods often lead to significant information loss and preclude the possibility of subsequent fine-grained scene and strategy analysis. *We advocate that the introduction of common sense and reasoning could pave the way towards ideal LO commentaries. On capturing atomic actions of game entities, using generic object detection techniques, the model can perform reasoning by: (1) constructing relational graphs among multiple entities, (2) conducting cross-temporal alignment with observation history, and (3) referring to game-specific knowledge and commonsense, thereby yielding fine-grained, comprehensive cognition of the current game state.*

Shallow Realization of SA. Current AI-GGC systems exhibit several critical limitations in SA. First, most systems lack long-horizon match memory, relying primarily on current actions and states for analysis, which makes it difficult to explain “Why?” or reliably anticipate “What’s next?”, both of which depend on broader historical context. *Maintaining memory of past actions and intermediate analytical states could help models form more coherent strategic interpretations.* Second, the modeling space of external augmentation remains underexplored. *Explainable AI methods provides a direct way to understand black-box game AIs, particularly interpretable game AI models [Li *et al.*, 2025a]. Meanwhile, deeper ToM modeling can help LLMs reason about players’ intentions and beliefs [Von Der Osten *et al.*, 2017].* Finally, SA is rarely internalized within LLMs, with most systems

relying on external analytical cues, constraining reasoning depth and consistency. *Aligning LLMs with strategic reasoning behaviors of game AIs at the parameter level is a promising direction for more robust analytical commentary.*

Lack of Cross-Capability Coordination. A common limitation of existing AI-GGC systems is the lack of coordination across core commentator capabilities. LO, SA, and HR are often treated as independent modules rather than as an integrated whole, with limited communicating or joint decision-making. Consequently, capabilities are optimized isolately, resulting in fragmented understanding and inconsistent commentary logic. While You [2025] adopts an end-to-end formulation, it primarily couples LO with commentary generation, leaving other core capabilities uncoordinated. *Cross-capability coordination can be promoted from both informational and functional perspectives. At the informational level, intermediate information produced by different capabilities can be partially shared. At functional level, SA and HR can impose fine-grained information demands on LO, specifying which details should be emphasized. Such demand-driven coordination remains unexplored in existing AI-GGC systems.*

Lack of a Standardized Benchmark and Reliable Evaluation. Despite the growing number of AI-GGC datasets spanning multiple modalities, game genres, and languages, the field still lacks a unified and standardized benchmark, making fair comparison across AI-GGC works difficult. Moreover, as discussed in Section 5, this problem is further exacerbated by the lack of reliable and convincing evaluation metrics. *A pressing direction is to develop a large-scale, multimodal AI-GGC benchmark with standardized data formats and fair evaluation. For evaluation, fine-tuning preference models from human labels is feasible but costly. In contrast, the taxonomy of commentary types proposed in our work enables more targeted LLM-as-a-Judge designs for different commentary types, alleviating challenges of multidimensionality and open-endedness. The relative importance and weights of evaluation dimensions can be examined through expert assessment or audience feedback [Valdivia and Burelli, 2025]. For analytical correctness, a promising direction is to decompose long reasoning chains into shorter, verifiable sub-chains and validate them progressively, reducing reliance on LLMs’ long-horizon reasoning evaluation.*

7 Conclusion

In this survey, we provide a systematic and comprehensive overview of AI-GGC. We propose a unified survey scheme that organizes existing work by game genres, core commentator capabilities, and commentary types, and use it to analyze methods, datasets, and evaluation metric. Moreover, we highlight key research gaps in current AI-GGC works and identify several promising directions for future work. We hope this survey offers a clear picture of the field and helps guide future research toward more mature AI commentators.

Contribution Statement

Qirui Zheng, Xingbo Wang, and Keyuan Cheng contributed equally to this work. Wenxin Li is the corresponding author.

References

- [Andrews *et al.*, 2024a] Peter Andrews, Oda Elise Nordberg, et al. Designing for automated sports commentary systems. In *IMX*, pages 75–93, 2024.
- [Andrews *et al.*, 2024b] Peter Andrews, Oda Elise Nordberg, et al. Aicommentator: A multimodal conversational agent for embedded visualization in football viewing. In *Proceedings of IUI*, pages 14–34, 2024.
- [Angeli *et al.*, 2010] Gabor Angeli, Percy Liang, et al. A simple domain-independent probabilistic approach to generation. In *Proceedings of EMNLP*, pages 502–512, 2010.
- [Chan *et al.*, 2022] Cheuk-Yiu Chan, Chun-Chuen Hui, et al. To start automatic commentary of soccer game with mixed spatial and temporal attention. In *Proceedings of TENCON*, pages 1–6, 2022.
- [Chen *et al.*, 2010] David L Chen, Joohyun Kim, et al. Training a multilingual sportscaster: Using perceptual context to learn language. *JAIR*, 37:397–435, 2010.
- [Chitrakala *et al.*, 2015] S Chitrakala, K Akshaya, et al. Story selection and recommendation system for colour commentary in cricket. 2015.
- [Cioppa *et al.*, 2022] Anthony Cioppa, Adrien Delière, et al. Scaling up soccernet with multi-view spatial localization and re-identification. *Scientific Data*, 9, 2022.
- [Delière *et al.*, 2021] Adrien Delière, Anthony Cioppa, et al. Soccernet-v2: A dataset and benchmarks for holistic understanding of broadcast soccer videos. In *CVPRW*, pages 4503–4514, 2021.
- [Denkowski and Lavie, 2014] Michael Denkowski and Alon Lavie. Meteor universal: Language specific translation evaluation for any target language. In *Proceedings of WMT*, pages 376–380, 2014.
- [Frith and Frith, 2005] Chris Frith and Uta Frith. Theory of mind. *Current biology*, 15(17):R644–R645, 2005.
- [Gautam *et al.*, 2024] Sushant Gautam, Mehdi Houshmand Sarkhoosh, et al. Soccernet-echoes: A soccer game audio commentary dataset. *ArXiv*, 2024.
- [Ghosh *et al.*, 2023] Debabrata Ghosh, Chiranjib Ghosh, et al. Automated cricket commentary generation using deep learning. *ICIASC*, 2023.
- [Giancola *et al.*, 2018] Silvio Giancola, Mohieddine Amine, et al. Soccernet: A scalable dataset for action spotting in soccer videos. In *CVPRW*, page 1792–179210, 2018.
- [Ishigaki *et al.*, 2023] Tatsuya Ishigaki, Goran Topić, et al. Audio commentary system for real-time racing game play. In *Proceedings of INLG: System Demonstrations*, 2023.
- [Jhamtani *et al.*, 2018] Harsh Jhamtani, Varun Gangal, et al. Learning to generate move-by-move commentary for chess games from large-scale social forum data. In *Proceedings of ACL*, pages 1661–1671, 2018.
- [Jung *et al.*, 2025] Sunghoon Jung, Hanmoe Kim, et al. Integrated ai system for real-time sports broadcasting: Player behavior, game event recognition, and generative ai commentary in basketball games. *Applied Sciences*, 2025.
- [Kameko *et al.*, 2015] Hirotaka Kameko, Shinsuke Mori, et al. Learning a game commentary generator with grounded move expressions. In *CIG*, 2015.
- [Kim and Choi, 2020] Byeong Jo Kim and Yong Suk Choi. Automatic baseball commentary generation using deep learning. In *Proceedings of SAC*, pages 1056–1065, 2020.
- [Kim *et al.*, 2025] Jaechang Kim, Jinmin Goh, et al. Bridging the gap between expert and language models: Concept-guided chess commentary generation and evaluation. In *Proceedings of NAACL-HLT*, 2025.
- [Kumano *et al.*, 2019] Tadashi Kumano, Manon Ichiki, et al. Generation of automated sports commentary from live sports data. In *BMSB*, pages 1–4, 2019.
- [Lee *et al.*, 2012] Greg Lee, Vadim Bulitko, et al. Sports commentary recommendation system (scores): machine learning for automated narrative. In *Proceedings of the 8th AAAI Conference on AIIDE*, page 32–37, 2012.
- [Lee *et al.*, 2014] Greg Lee, Vadim Bulitko, et al. Automated story selection for color commentary in sports. *IEEE T-CIAIG*, 6(2):144–155, 2014.
- [Lee *et al.*, 2022] Andrew Lee, David Wu, et al. Improving chess commentaries by combining language models with symbolic reasoning engines. *ArXiv*, 2022.
- [Li *et al.*, 2025a] Lingfeng Li, Yunlong Lu, et al. Mxplainer: Explain and learn insights by imitating mahjong agents. *Algorithms*, 18(12):738, 2025.
- [Li *et al.*, 2025b] Yunxin Li, Zhenyu Liu, et al. Perception, reason, think, and plan: A survey on large multimodal reasoning models. *ArXiv*, 2025.
- [Liang *et al.*, 2009] Percy Liang, Michael I Jordan, et al. Learning semantic correspondences with less supervision. In *Proceedings of ACL-IJCNLP*, pages 91–99, 2009.
- [Lin, 2004] Chin-Yew Lin. ROUGE: A package for automatic evaluation of summaries. In *Text Summarization Branches Out*, pages 74–81, 2004.
- [Liu *et al.*, 2023] Yang Liu, Dan Iter, et al. G-eval: Nlg evaluation using gpt-4 with better human alignment. In *Proceedings of EMNLP*, pages 2511–2522, 2023.
- [Ma *et al.*, 2025] Feihu Ma, Xuechen Liang, et al. Strategies for enhancing card game analysis: Exploring large language models in imperfect information settings. *IEEE Access*, 2025.
- [Matsuyoshi *et al.*, 2018] Suguru Matsuyoshi, Hirotaka Kameko, et al. Annotating modality expressions and event factuality for a japanese chess commentary corpus. In *Proceedings of LREC*, 2018.
- [Mei *et al.*, 2016] Hongyuan Mei, Mohit Bansal, et al. What to talk about and how? selective generation using LSTMs with coarse-to-fine alignment. In *Proceedings of NAACL-HLT*, 2016.
- [Mkhallati *et al.*, 2023] Hassan Mkhallati, Anthony Cioppa, et al. Soccernet-caption: Dense video captioning for soccer broadcasts commentaries. *CVPRW*, 2023.

- [Mori *et al.*, 2016] Shinsuke Mori, John Richardson, et al. A Japanese chess commentary corpus. In *LREC*, 2016.
- [Papineni *et al.*, 2002] Kishore Papineni, Salim Roukos, et al. Bleu: a method for automatic evaluation of machine translation. In *ACL*, pages 311–318, 2002.
- [Qi *et al.*, 2023] Ji Qi, Jifan Yu, et al. Goal: A challenging knowledge-grounded video captioning benchmark for real-time soccer commentary generation. In *Proceedings of CIKM*, pages 5391–5395, 2023.
- [Ranella and Eger, 2023] Noah Ranella and Markus Eger. Towards automated video game commentary using generative ai. In *EXAG@AIIDE*, 2023.
- [Rao *et al.*, 2024] Jiayuan Rao, Haoning Wu, et al. Matchtime: Towards automatic soccer game commentary generation. In *Proceedings of EMNLP*, 2024.
- [Rauf *et al.*, 2020] Muhammad Arslan Rauf, Haseeb Ahmad, et al. Extraction of strong and weak regions of cricket batsman through text-commentary analysis. In *INMIC*, pages 1–6, 2020.
- [Reiter and Belz, 2009] Ehud Reiter and Anja Belz. An investigation into the validity of some metrics for automatically evaluating natural language generation systems. *Computational Linguistics*, 35(4):529–558, 2009.
- [Ringer *et al.*, 2019] Charles Ringer, James Alfred Walker, et al. Multimodal joint emotion and game context recognition in league of legends livestreams. In *IEEE CoG*, 2019.
- [Shi *et al.*, 2024] Lin Shi, Chiyu Ma, et al. Judging the judges: A systematic study of position bias in llm-as-a-judge. *arXiv preprint*, 2024.
- [Silver *et al.*, 2016] David Silver, Aja Huang, et al. Mastering the game of go with deep neural networks and tree search. *nature*, 529(7587):484–489, 2016.
- [Siu *et al.*, 2023] Wan-Chi Siu, H Anthony Chan, et al. On the completion of automatic football game commentary system with deep learning. In *IWAIT*, 2023.
- [Tanaka and Simo-Serra, 2021] Tsunehiko Tanaka and Edgar Simo-Serra. Lol-v2t: Large-scale esports video description dataset. In *CVPRW*, pages 4552–4561, 2021.
- [Taniguchi *et al.*, 2019] Yasufumi Taniguchi, Yukun Feng, et al. Generating live soccer-match commentary from play data. In *AAAI*, 2019.
- [Tomlin *et al.*, 2022] Nicholas Tomlin, Andre He, et al. Understanding game-playing agents with natural language annotations. In *Proceedings of ACL*, pages 797–807, 2022.
- [Ul Abideen *et al.*, 2021] Zain Ul Abideen, Saira Jabeen, et al. Ball-by-ball cricket commentary generation using stateful sequence-to-sequence model. In *ComTech*, 2021.
- [Valdivia and Burelli, 2025] Arturo Valdivia and Paolo Burelli. Evaluating quality of gaming narratives co-created with ai. In *CoG*, 2025.
- [Vijayakumar *et al.*, 2025] Adarsh Vijayakumar, Amal Toms, et al. Player tracking-integrated soccer game commentary generation. *IJSAT*, 16(2), 2025.
- [Von Der Osten *et al.*, 2017] Friedrich Burkhard Von Der Osten, Michael Kirley, et al. The minds of many: Opponent modeling in a stochastic game. In *IJCAI*, 2017.
- [Wang and Yoshinaga, 2021] Zihan Wang and Naoki Yoshinaga. From esports data to game commentary: Datasets, models, and evaluation metrics. In *DEIM*, 2021.
- [Wang *et al.*, 2025] Xinyu Wang, Bohan Zhuang, et al. Are large vision language models good game players? *arXiv preprint*, 2025.
- [Xu and Li, 2007] Jun Xu and Hang Li. Adarank: a boosting algorithm for information retrieval. *ACM*, 2007.
- [Xu *et al.*, 2021] Junjie H. Xu, Zhou Fang, et al. Fighting game commentator with pitch and loudness adjustment utilizing highlight cues. In *GCCE*, pages 366–370, 2021.
- [Xu *et al.*, 2023] Junjie H. Xu, Yu Nakano, et al. Cs-lol: a dataset of viewer comment with scene in e-sports live-streaming. In *Proceedings of CHIIR*, page 422–426, 2023.
- [Yan *et al.*, 2016] Fei Yan, Krystian Mikolajczyk, et al. Generating commentaries for tennis videos. In *ICPR*, 2016.
- [Yan *et al.*, 2022] Yichao Yan, Ning Zhuang, et al. Fine-grained video captioning via graph-based multi-granularity interaction learning. *IEEE TPAMI*, 2022.
- [Ye *et al.*, 2024] Jiayi Ye, Yanbo Wang, et al. Justice or prejudice? quantifying biases in llm-as-a-judge. *arXiv preprint*, 2024.
- [You *et al.*, 2025] Ling You, Wenxuan Huang, et al. Time-soccer: An end-to-end multimodal large language model for soccer commentary generation. In *ACM MM*, 2025.
- [Yu *et al.*, 2018] Huanyu Yu, Shuo Cheng, et al. Fine-grained video captioning for sports narrative. In *CVPR*, pages 6006–6015, 2018.
- [Zang *et al.*, 2019] Hongyu Zang, Zhiwei Yu, et al. Automated chess commentator powered by neural chess engine. In *Proceedings of ACL*, pages 5952–5961, 2019.
- [Zhang and Eickhoff, 2021] Ruochen Zhang and Carsten Eickhoff. Soccer: An information-sparse discourse state tracking collection in the sports commentary domain. In *Proceedings of NAACL*, 2021.
- [Zhang *et al.*, 2021] Yiming Zhang, Albertus Agung, et al. An audience participation game with difficulty adjustment and rap-style commentary based on audience inputs. In *GCCE*, pages 845–846, 2021.
- [Zhang *et al.*, 2022] Dawei Zhang, Sixing Wu, et al. MOBA-E2C: Generating MOBA game commentaries via capturing highlight events from the meta-data. In *Findings of the EMNLP*, pages 4545–4556, 2022.
- [Zhang *et al.*, 2024a] Benhui Zhang, Junyu Gao, et al. A descriptive basketball highlight dataset for automatic commentary generation. In *ACM MM*, 2024.
- [Zhang *et al.*, 2024b] Zhihao Zhang, Feiqi Cao, et al. Game-mug: Multimodal oriented game situation understanding and commentary generation dataset. *ArXiv*, 2024.