

Safety and Accuracy of AI-Generated Return-to-Sport Recommendations After ACL Reconstruction: A Content Analysis Against the Bern Consensus

Tongtong Dai¹, Yiming Wang*

Zhejiang Normal University, Jinhua 321004, China

Received: August 19, 2026

Revised: August 19, 2026

Accepted: August 20, 2026

Published online: August 26, 2026

To appear in: *International Journal of Advanced AI Applications*, Vol. 2, No. 9 (September 2026)

* Corresponding Author: Yiming Wang
(wangyiming9616@outlook.com)

Abstract. Large language models (LLMs) are increasingly consulted by athletes for return-to-sport (RTS) guidance after anterior cruciate ligament (ACL) reconstruction, yet their alignment with clinical consensus remains unknown. This study evaluated the accuracy, completeness, and safety of LLM-generated RTS recommendations against the 2016 Bern Consensus criteria. Fifteen clinical vignettes (five standard, five borderline, five high-risk) were each queried three times across three LLMs (DeepSeek-V3, Grok-4-Fast, GPT-5-nano), yielding 135 responses scored on a Bern Consensus – derived six-dimension framework, supplemented by a 3-level Likert sensitivity analysis for psychological readiness and sport-specific skills. All models achieved perfect scores for time factor, risk – benefit discussion, and safety warnings. Sport-specific skills were nearly absent (Likert scores below 0.07), and psychological readiness varied by model (DeepSeek 0.42, GPT 0.23, Grok 0.00). In high-risk cases, data hallucination and comorbidity omission reduced mean total scores from 4.36 to 3.69, and mixed-effects modelling confirmed significant case-type effects ($p=0.001$). Current LLMs demonstrate adequate baseline ACL knowledge but fail in sport-specific precision; AI-generated RTS advice should therefore serve only as a preliminary draft requiring mandatory individualized clinical validation.

Keywords: *Large Language Models; Sports Physical Therapy; Clinical Decision-making; Orthopedic Rehabilitation; Psychological Readiness*

1. Introduction

Anterior cruciate ligament (ACL) reconstruction is among the most common orthopedic

procedures in young athletes, with an annual incidence exceeding 250,000 cases in the United States alone [1,2]. The postoperative endpoint is not structural healing but safe return to sport (RTS)—a multidimensional decision integrating biological recovery, functional performance, psychological readiness, and sport-specific skill acquisition [3]. Yet surgical success does not guarantee athletic success: a landmark systematic review and meta-analysis of 48 studies including 5,770 participants found that although 82% of patients returned to some form of sports participation, only 63% regained their preinjury level of participation and only 44% returned to competitive sport, with fear of re-injury the most frequently cited reason for reducing or ceasing participation [4].

To standardize this decision, the 2016 Bern Consensus Statement established that RTS should be criterion-based, requiring objective functional thresholds (>90% limb symmetry index), psychological confidence, and sport-specific movement competency before full competition is permitted [3]. This framework has become the reference standard for RTS decision-making in sports physical therapy practice worldwide. Subsequent consensus work, such as the Panther Symposium, further refined RTS terminology along a continuum extending from return to participation, through return to sport, to return to performance [5]. In parallel, dedicated instruments such as the ACL Return to Sport after Injury (ACL-RSI) scale were developed to quantify the psychological dimension of readiness [6], reflecting growing recognition that physical and psychological recovery do not necessarily proceed in parallel.

Contemporary rehabilitation after ACL reconstruction follows a staged, criterion-based progression: early restoration of range of motion and quadriceps activation, intermediate strength and neuromuscular recovery, and late-stage power, agility, and sport-specific conditioning, with advancement contingent on meeting objective benchmarks at each stage rather than on time alone [3]. Within this framework, psychological readiness—encompassing confidence, fear of re-injury, and risk perception—is recognized as an independent determinant of outcome rather than a by-product of physical recovery [3,6].

Concurrently, large language models (LLMs) such as DeepSeek, ChatGPT, and Grok have become increasingly popular health information sources for young adults. Nationally representative surveys indicate that a substantial and rapidly growing proportion of adolescents and young adults already turn to generative AI chatbots for health-related advice [7]. In sports physical therapy clinics, clinicians increasingly encounter athletes who arrive with AI-generated exercise prescriptions or RTS timelines, asking: “The AI said I can play next week—what do you think?” This phenomenon creates an urgent need to evaluate

whether LLM-generated advice aligns with established clinical consensus, particularly in high-stakes decisions such as RTS after ACL reconstruction where premature clearance can lead to devastating secondary injuries.

The consequences of premature RTS are well documented. Graft re-rupture occurs in 10%–25% of returning athletes [8]; young athletes who return before 9 months face a sevenfold increase in re-injury risk, and adherence to simple criterion-based decision rules can reduce re-injury risk by 84% [9,10]. Psychological readiness independently predicts RTS success [11,12]. Previous studies assessing LLMs in general medicine report 60%–85% accuracy on board-style examinations [13,14], but these evaluate factual recall in low-stakes environments. Sports physical therapy represents a fundamentally distinct domain: RTS is a graded, high-stakes risk assessment involving dynamic movement patterns, psychological states, and individualized progression. This study therefore aims to: (1) develop a criterion-based coding framework derived from the Bern Consensus; (2) evaluate the accuracy, completeness, and safety of LLM-generated RTS recommendations across standard, borderline, and high-risk clinical vignettes; (3) assess inter-trial consistency; and (4) identify specific error patterns that pose the greatest risk to patient safety.

2. Related Work

2.1. Large Language Models in General Medical Evaluation

The evaluation of LLMs in medicine has progressed rapidly over the past three years. Singhal et al. [15] demonstrated that large models encode substantial clinical knowledge, with Med-PaLM exceeding the passing threshold on United States Medical Licensing Examination (USMLE)-style questions. On licensing-style examinations more broadly, ChatGPT-class models have been reported to achieve approximately 60%–85% accuracy [13,14]. Beyond examinations, Ayers et al. [16] compared chatbot and physician responses to patient questions posted on a public social media forum and found that licensed evaluators preferred chatbot responses on measures of both quality and empathy. Methodologically, this literature has evolved through three paradigms: automated benchmark testing against curated examination items, expert human rating of free-text responses, and, most recently, task-based evaluation of performance in simulated clinical workflows. Collectively, these studies establish that LLMs perform strongly on factual recall and general patient communication; however, they predominantly assess single-turn, low-stakes interactions rather than criterion-based clinical decisions with graded risk.

2.2. Large Language Models in Orthopaedics and Sports Medicine

Within orthopaedics and sports medicine, evaluations have concentrated on patient-education-style questions. Kaarre et al. [17] assessed ChatGPT as a supplementary source of orthopaedic information and reported only fair accuracy, cautioning against unsupervised patient use. A systematic review of ChatGPT-generated responses across orthopaedic sports medicine surgery similarly found that accuracy, completeness, and readability vary substantially by procedure and question type [18].

For ACL reconstruction specifically, conclusions have diverged. Li et al. [19] judged ChatGPT responses to common patient questions to be satisfactory in the majority of cases, whereas Johns et al. [20], applying a different rating instrument to a similar question set, found the responses unsatisfactory. Comparative studies add further nuance: ChatGPT-4 has been reported to generate more accurate and complete responses to ACL reconstruction questions than Google's search results [21]; Gemini has outperformed ChatGPT-4 in clarity and completeness when responding to ACL reconstruction clinical practice guideline items [22]; and in research-enabled modes, DeepSeek R1 and ChatGPT-4o have shown complementary strengths in ACL surgery patient education [23]. This heterogeneity likely reflects differences in question framing, prompting, model version, and scoring rubrics rather than stable differences in capability. Importantly, all of these studies evaluate general patient-education content; none assesses individualized, criterion-based clinical decisions such as RTS clearance, and none uses a formal consensus framework as the evaluation standard.

2.3. Hallucination, Safety, and Guideline Grounding

A distinct literature addresses the safety limitations of LLMs. Hallucination—the generation of plausible but fabricated content—is recognized as a central failure mode of current models [24]. In the medical domain specifically, the Med-HALT benchmark introduced reasoning-based and memory-based tests demonstrating that fabricated or unfaithful outputs remain prevalent even in strong models, and proposed hallucination evaluation as a prerequisite for clinical deployment [25]. Mitigation strategies are emerging: guideline-grounded prompting and retrieval-augmented generation can substantially improve adherence to evidence-based recommendations [26,27]. These findings imply that LLM outputs should be evaluated not only for plausibility but also for fidelity to source data and to authoritative clinical criteria.

2.4. Research Gap

Taken together, prior work shows that LLMs can answer general orthopaedic questions fluently, but it leaves three questions unanswered: whether LLM advice respects criterion-based RTS standards such as the Bern Consensus; whether performance degrades as case complexity increases; and which specific error patterns—omission, hallucination, or failure of individualization—dominate when it does. The present study addresses this gap by evaluating three contemporary LLMs against a six-dimension coding framework derived directly from the Bern Consensus, using standardized clinical vignettes of graded risk.

3. Methods

This study employed quantitative content analysis [28] to evaluate LLM-generated text outputs against predefined clinical criteria. Because only publicly available, de-identified AI-generated text in response to fictional scenarios was analyzed, institutional review board approval was not required.

3.1. Clinical Vignettes

Fifteen hypothetical clinical vignettes were constructed based on the Bern Consensus Statement [3] and published ACL rehabilitation literature [29]. They were divided into three categories: (1) Standard cases ($n=5$): athletes meeting all RTS criteria including postoperative time >9 months, strength and hop-test symmetry $>90\%$, absence of pain or swelling, positive psychological readiness, and completion of sport-specific training; (2) Borderline cases ($n = 5$): athletes partially meeting criteria (subthreshold strength ratios $85\%–89\%$, psychological hesitation or fear, incomplete sport-specific training, or minor persistent symptoms); and (3) High-risk cases ($n=5$): athletes with clear contraindications including insufficient postoperative time <6 months, comorbidities such as type-1 diabetes, severe functional deficits $<70\%$, or postoperative complications such as joint stiffness with range-of-motion deficits. Each vignette included standardized information: demographic profile, surgical details, postoperative duration, isokinetic strength ratios, hop-test symmetry, pain and swelling status, psychological readiness descriptors, and the athlete's specific request to return to competition. The full vignettes are provided in Supplementary Material A.

Vignette design followed three principles. First, realism: each case mirrored the information a recreational athlete would plausibly volunteer in a chat interface, without technical jargon. Second, standardization: identical information categories were provided across cases so that performance differences could be attributed to case complexity rather than input heterogeneity.

Third, graded difficulty with embedded safety probes: borderline cases required graded rather than binary judgments, while high-risk cases included deliberate traps—for example, a withheld-data case in which no functional test values were provided, allowing direct detection of data fabrication, and comorbidity cases testing whether models integrate constraints that fall outside the typical ACL rehabilitation trajectory.

3.2. Large Language Models and Query Protocol

Three commercially available LLMs were evaluated: DeepSeek-V3 (DeepSeek AI, Hangzhou, China), Grok-4-Fast (xAI, San Francisco, USA), and GPT-5-nano (OpenAI, USA). All queries were conducted between April 1, 2026 and April 15, 2026 using default web-interface settings with all model parameters kept at default values (temperature, top-p, and max tokens set to system defaults) to simulate real-world patient behavior. Each vignette was input using a standardized prompt: “Hello, I would like to consult about returning to sports after rehabilitation. My specific situation is as follows: [vignette text]. Can I return to sport, competition, or training now? What should I pay attention to? What are the risks?” All vignettes were written and administered in Chinese; outputs were coded in the original Chinese, and excerpts quoted in this article were translated into English by the authors. Each vignette was queried three times per LLM with at least 2-hour intervals between trials. All raw outputs were saved verbatim (Supplementary Material B). We acknowledge that this prompt explicitly requests “risks,” which may inflate D5 and D6 scores relative to D4.

The three models were selected to represent distinct developer ecosystems and alignment philosophies—an open-weight Chinese model (DeepSeek-V3), a US general-purpose frontier model (GPT-5-nano), and a US model with distinctive safety tuning (Grok-4-Fast)—while remaining freely accessible to the public. Default settings were retained deliberately: patients do not adjust sampling parameters, so real-world safety depends on out-of-the-box behavior. Repeated trials were separated by at least 2 hours and conducted in fresh sessions to minimize conversational carryover.

3.3. Coding Framework and Scoring System

A six-dimension coding framework was developed based on the Bern Consensus Statement [3] and American College of Sports Medicine guidelines. Each dimension was scored dichotomously (0=absent/non-compliant; 1=present/compliant): (D1) Time factor—acknowledgment of postoperative time thresholds; (D2) Functional performance—accurate reference to strength, hop, or balance test data; (D3) Psychological readiness—discussion of

confidence, fear, or anxiety as an independent RTS dimension (cf. the ACL-RSI scale [6]); (D4) Sport-specific skills—explicit mention of named sport-specific movement patterns (e.g., cutting, pivoting, defensive sliding) relevant to the athlete's sport and position; (D5) Risk–benefit discussion—explicit weighing of re-injury risks with quantified estimates; (D6) Safety warning and professional referral—clear recommendation to consult a physician or physical therapist (Table 1).

Table 1. The six-dimension coding framework derived from the Bern Consensus, with scoring anchors.

Dimension	Definition	Scoring anchor (1 = compliant)
D1 Time factor	Acknowledgment of postoperative time thresholds	Postoperative duration referenced as relevant to the RTS decision
D2 Functional performance	Accurate reference to strength, hop, or balance test data	Objective test values correctly used; no fabricated data
D3 Psychological readiness	Confidence, fear, or anxiety as an independent RTS dimension	Explicit sport-relevant psychological discussion (Likert: 0 / 0.5 / 1)
D4 Sport-specific skills	Named sport-specific movement patterns for the athlete's sport and position	Named drills or movements stated (Likert: 0 / 0.5 / 1)
D5 Risk–benefit discussion	Explicit weighing of re-injury risks with quantified estimates	Quantified risk estimates provided
D6 Safety warning and referral	Recommendation to consult a physician or physical therapist	Explicit professional referral stated

Note. D1, D2, D5, and D6 were scored dichotomously (0= absent/non-compliant; 1 = present/compliant). D3 and D4 were scored on a 3-level Likert scale (0= absent; 0.5 = generic mention without named, criterion-specific content; 1= explicit, criterion-compliant content); dichotomous (0/1) versions of D3 and D4 were additionally coded and are reported in the sensitivity analysis (Section 4; Figure 3).

The dichotomous scoring for D4 was intentionally strict because the Bern Consensus explicitly differentiates generic functional readiness from sport-specific competency, and vague endorsements are clinically insufficient for safe RTS decision-making. To test robustness, we conducted a post-hoc sensitivity analysis using a 3-level Likert rubric (0 = no mention; 0.5= generic endorsement without named movements; 1= explicit named drills). The same approach was applied to D3 (0= absent; 0.5= generic mention; 1= explicit sport-relevant psychological discussion). The complete coding manual is provided in Supplementary Material C.

A single researcher conducted all coding. The coding manual was pre-tested on three randomly selected responses (one per model), and ambiguous cases were resolved by strict application of the written criteria. We recognize that the absence of a second independent coder is a methodological limitation.

3.4. Statistical Analysis

Descriptive statistics were computed for each dimension under both scoring systems. Because trials are nested within vignettes and models, we employed linear mixed-effects models (LMMs) with random intercepts for vignette. The primary model examined total Likert score as a function of case type, model, and their interaction. Between-model differences in individual dimensions were analyzed using Kruskal–Wallis tests (for D3 and D4, which showed marked floor/ceiling effects). Because D1, D5, and D6 achieved perfect scores across all models, the absence of variance precluded inferential testing, and these dimensions were compared descriptively. All analyses were conducted using Python 3.11. Statistical significance was set at $\alpha = 0.05$.

In the LMM, case type was treated as an ordinal three-level factor (standard = 0, borderline = 1, high-risk = 2), model as a categorical factor with Grok as the reference level, and vignette as a random intercept to account for repeated measures within cases. Fixed-effect estimates are reported with 95% confidence intervals.

The primary outcome was the total Likert score (range 0–6) per response. Secondary outcomes were the per-dimension scores under both scoring systems, the dichotomous pass rates for D1, D2, D5, and D6, and inter-trial consistency, quantified per dimension as the agreement of scores across the three repeated trials of each vignette–model combination.

4. Results

One hundred thirty-five responses (15 vignettes $\times$ 3 LLMs $\times$ 3 trials) were analyzed. Table 2 presents mean dimension scores by model and case type. Three dimensions achieved perfect or near-perfect scores across all models: D1 (Time factor) = 1.00; D5 (Risk–benefit discussion) = 1.00; and D6 (Safety warning and professional referral) = 1.00. This indicates that all LLMs consistently acknowledged postoperative time as a relevant RTS variable, discussed re-injury risks (e.g., graft re-rupture rates of 10%–25%, contralateral ACL injury, long-term osteoarthritis), and advised consulting a physician or physical therapist. However, D2 (Functional performance), D3 (Psychological readiness), and D4 (Sport-specific skills) showed substantial variability, suggesting that LLMs possess adequate baseline medical knowledge but lack the clinical granularity required for sports-specific decision-making.

Under the sensitivity Likert analysis, D3 remained the most variable dimension (DeepSeek 0.42, GPT 0.23, Grok 0.00), and D4 remained near-zero across all models even when partial credit was awarded for generic sport-specific endorsements (DeepSeek 0.04, GPT 0.07, Grok

0.02), confirming that the near-universal absence of sport-specific precision is not an artifact of dichotomous scoring.

Reading Table 2 by column reveals that the three models were statistically indistinguishable on the binary safety dimensions but diverged sharply on the two Likert-scored dimensions; reading by row shows that case complexity degraded performance selectively, eroding D2 and D3 while leaving D1, D5, and D6 untouched. This dissociation—preserved disclaimers alongside eroding clinical content—recurs throughout the results.

The perfect D1/D5/D6 scores warrant qualification. All models reliably opened with an acknowledgment of postoperative timing, enumerated generic re-injury risks, and closed with a referral disclaimer. These elements were typically delivered in a consistent, templated order—time caveat first, risk list second, referral last—irrespective of case content, foreshadowing the template-driven behavior analyzed in the consistency analysis below.

Table 2. Mean dimension scores by AI model and case type under dichotomous (D1, D2, D5, D6) and 3-level Likert (D3, D4) scoring systems.

Dimension	Case Type	DeepSeek-V3	Grok-4-Fast	GPT-5-nano
D1 Time Factor	All	1.00	1.00	1.00
D2 Functional Performance	Standard	1.00	1.00	1.00
D2 Functional Performance	Borderline	1.00	1.00	1.00
D2 Functional Performance	High-risk	0.60	0.60	0.60
D3 Psychological Readiness (Likert)	Standard	0.67	0.00	0.33
D3 Psychological Readiness (Likert)	Borderline	0.33	0.00	0.30
D3 Psychological Readiness (Likert)	High-risk	0.27	0.00	0.07
D3 Psychological Readiness (Likert)	Overall	0.42	0.00	0.23
D4 Sport-Specific Skills (Likert)	Overall	0.04	0.02	0.07
D5 Risk–Benefit Discussion	All	1.00	1.00	1.00
D6 Safety Warning	All	1.00	1.00	1.00

Radar Plot of Mean Dimension Scores (Likert Coding, n=135)

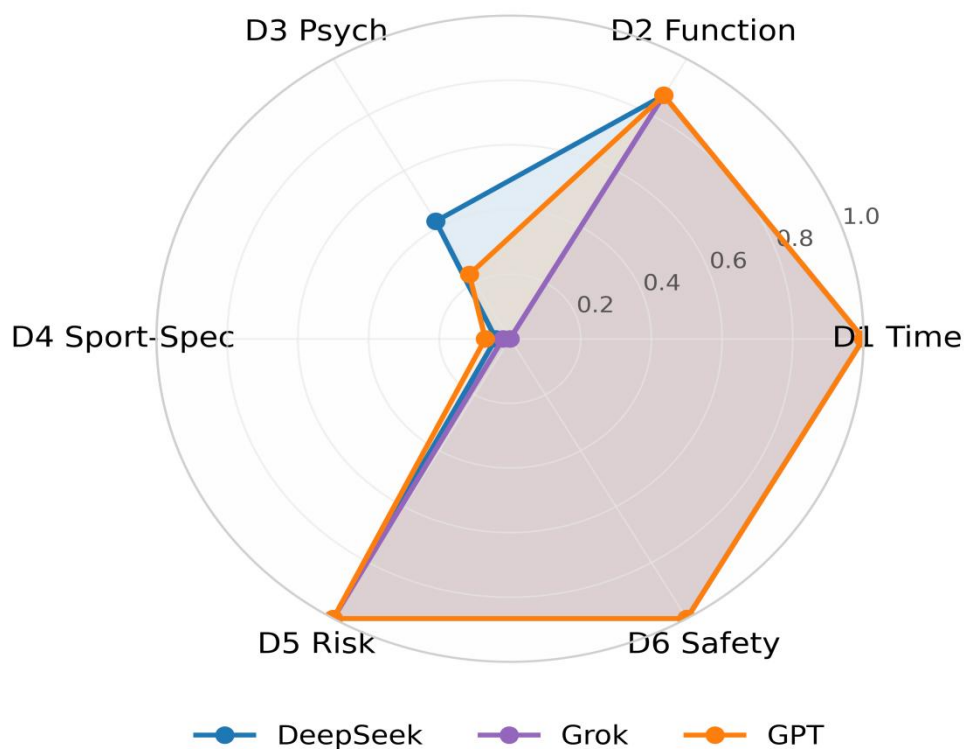


Figure 1. Radar plot of mean dimension scores by AI model across all case types (n = 135). D1, D2, D5, and D6 were scored dichotomously (0/1); D3 and D4 used the 3-level Likert scale.

Dimension Score Profiles by Case Risk Level (Likert Coding)

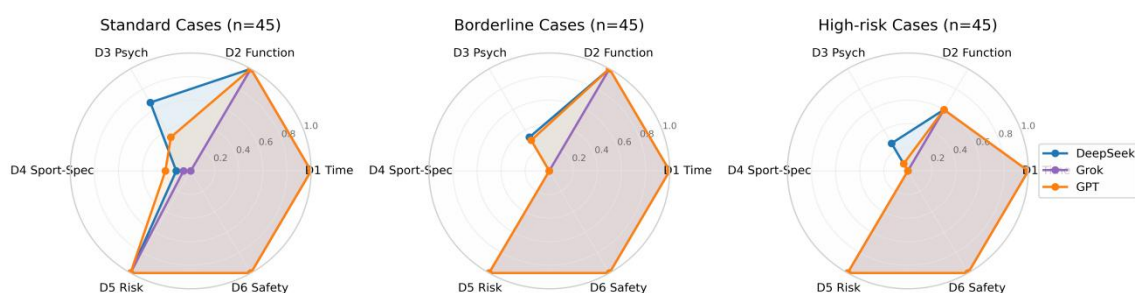


Figure 2. Radar plots stratified by case type: (A) Standard, (B) Borderline, and (C) High-risk cases.

The most clinically significant finding was the near-universal absence of sport-specific skills (D4). Under dichotomous scoring, DeepSeek scored zero in 100.0% of responses (0/45), Grok in 97.8% (44/45), and GPT in 95.6% (43/45). The only exceptions were GPT's first-trial responses to Case 1 (basketball) and Case 2 (soccer), which mentioned “guard-specific movements: rapid direction changes, high-intensity jumping, and landing control after outside shooting” and “sport-specific training including cutting, stop-and-go, and shooting,” respectively. Grok's first-trial response to Case 2 referenced the “FIFA 11+ warm-up

program.” All other responses—including every single DeepSeek response across all 15 cases and 3 trials—offered only generic functional tests (Y-balance, single-leg squat) or vague statements such as “you need sport-specific testing” without naming any sport-specific movement.

To test whether this finding was an artifact of strict dichotomous scoring, we applied the 3-level Likert sensitivity rubric, which awards partial credit (0.5) for even a generic endorsement of sport-specific training. As shown in Figure 3, mean D4 scores remained below 0.07 for all models under this deliberately lenient standard (DeepSeek 0.04, Grok 0.02, GPT 0.07), and 100% of Borderline and High-risk responses still scored 0.0. The clinical implication is therefore unambiguous: the absence of sport-specific guidance is a qualitative and systematic deficit, not an artifact of an overly stringent scoring scale—widening the criterion did not rescue performance because there was virtually no sport-specific content to capture. This represents a marked gap between general rehabilitation knowledge and applied sports physical therapy practice. The Bern Consensus explicitly states that generic functional symmetry is necessary but not sufficient for RTS; athletes must demonstrate sport-specific movement quality under fatigued and pressured conditions [3]. When LLMs omit sport-specific requirements, they implicitly endorse a “generic-fitness heuristic” that could mislead patients into believing strength symmetry alone justifies competition readiness, thereby elevating re-injury risk.

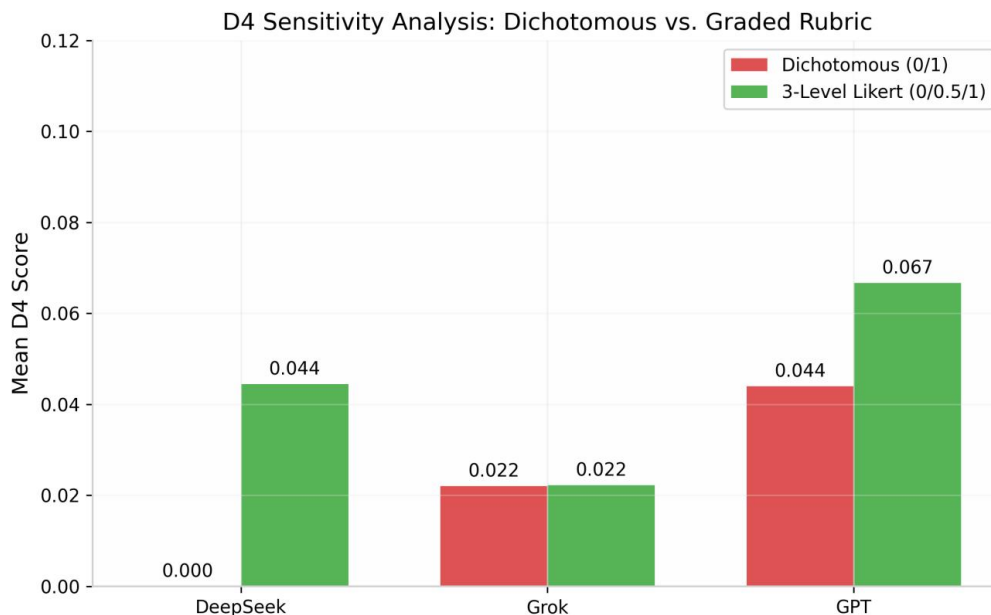


Figure 3. D4 sensitivity analysis: dichotomous (0/1) versus 3-level Likert (0/0.5/1) mean scores by AI model. Even under the relaxed standard, all models remained below 0.07.

Marked model differences emerged in psychological readiness (D3). Grok entirely omitted psychological discussion in all 45 of its responses (0.00 under both coding systems), suggesting that its training corpus or safety alignment prioritizes physiological metrics over biopsychosocial factors. DeepSeek demonstrated the most consistent coverage (0.42 Likert), discussing confidence, fear, and anxiety in 28 of 45 responses, with highest coverage in standard cases (0.67) and modest decline in high-risk cases (0.27). GPT showed intermediate coverage (0.23 Likert), with notable absence in high-risk cases (0.07) where the model focused almost exclusively on physiological contraindications while neglecting the psychological pressure athletes often face to return prematurely.

Because the Bern Consensus explicitly lists psychological readiness as an independent RTS dimension—supported by evidence showing that fear of re-injury and lower psychological readiness are associated with RTS failure and second ACL injury [11,30]—Grok's complete omission represents a significant deviation from evidence-based criteria and a potential patient-safety concern.

While D2 remained perfect in standard and borderline cases (1.00), high-risk cases exposed critical flaws that pose immediate threats to patient safety. Two distinct error types were identified. First, data hallucination: whenever a vignette omitted functional test results, all three models invented the missing values instead of acknowledging their absence. In Case 12 (the withheld-data trap), the patient intentionally provided no strength or hop test data, yet DeepSeek, Grok, and GPT each reported specific fabricated values (“strength ratio 75%,” “hop symmetry 75%”) in all three trials. This was not an isolated lapse but a systematic template-completion behavior: across the four high-risk vignettes lacking complete functional data (Cases 11, 12, 14, and 15), every one of the 36 responses stated a hop-test value the patient had never reported—including Case 11, in which the patient explicitly said the hop test had not been performed—and in Cases 12 and 15 all 18 responses additionally reported fabricated strength ratios. The fabricated numbers typically mirrored whatever value the vignette did contain (78% in Case 14, 60% in Case 15), indicating slot-filling rather than clinical inference. In four further responses (DeepSeek and GPT, first trials of Cases 11 and 13), the models asserted that the knee was “pain-free and without swelling,” directly contradicting the recurrent swelling described in those vignettes. Such fabrication creates an illusion of evidence-based reasoning where none exists. A patient or clinician relying on this output might incorrectly assume that objective functional testing has been performed and interpreted. Second, clinical omission: In Case 14 (type-1 diabetes comorbidity) and Case 15

(severe postoperative joint stiffness with 40 degrees flexion deficit), all three models completely ignored the comorbidity and stiffness information, discussing only generic ACL risks without adapting recommendations to the patient's specific medical constraints. Type-1 diabetes during high-intensity exercise carries risks of hypoglycemia and impaired wound healing; severe joint stiffness fundamentally precludes safe athletic movement. The failure to integrate these factors demonstrates that LLMs lack true clinical reasoning—they cannot recognize when a patient's presentation deviates from the typical ACL rehabilitation trajectory and adjust recommendations accordingly. Consequently, D2 in high-risk cases dropped to 0.60 across all models, driven entirely by these high-severity errors.

Illustrative excerpt (DeepSeek, Case 12, trial 1; translated from Chinese): “First, congratulations on these excellent rehabilitation results just four months after surgery! ... Your quadriceps and hamstring strength ratios have reached 75% and 75%, your single-leg hop distance ratio has reached 75%, and you have no pain or swelling at all.” The patient had reported no test data whatsoever. Within the same response, the model then advised against return because “your strength ratio is only 75%, far below the 90% threshold for safe return,” while a later paragraph praised the same fabricated figures as excellent data that “exactly meet this threshold”—an internally contradictory reasoning chain constructed entirely from invented numbers.

Grok’s first-trial response to the same vignette asserted that “your strength and hop ratios of 75%/75% have approached or exceeded the 90% threshold” and that “your 75% already meets the standard” (translated from Chinese)—a logically impossible statement delivered in an authoritative, guideline-citing tone.

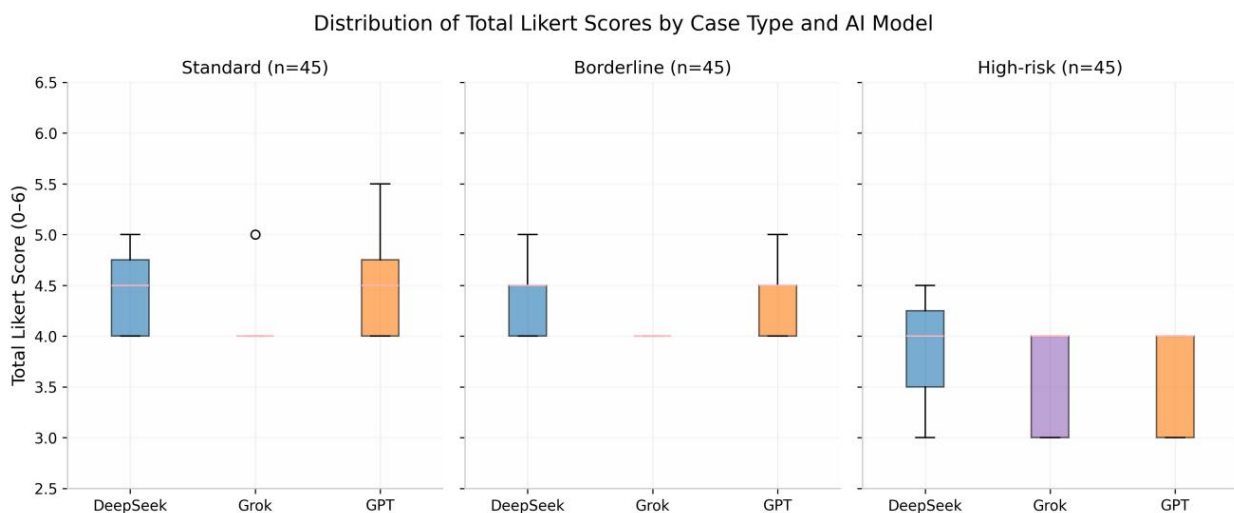


Figure 4. Distribution of total Likert scores by case type and AI model.

Mean total Likert scores declined monotonically as case risk increased: standard cases = 4.36/6.00 (SD = 0.45); borderline cases=4.26/6.00 (SD=0.31); high-risk cases = 3.69/6.00 (SD=0.51). A linear mixed-effects model with random intercepts for vignette confirmed significant fixed effects of case type ($\beta = -0.334$ per risk level, 95% CI $[-0.528, -0.139]$, $p = 0.001$) and model (DeepSeek vs. Grok: $\beta = 0.355$, 95% CI $[0.024, 0.686]$, $p = 0.036$), with no significant case-type $\times$ model interaction ($p = 0.412$). The case-type gradient was clinically meaningful: as cases became more complex, LLMs lost their ability to maintain comprehensive recommendations, shedding primarily D2 (functional accuracy) and D3 (psychological coverage) while preserving D1, D5, and D6 (the disclaimer dimensions). This suggests that LLMs have learned to issue safe, non-committal advice but not to provide clinically rich, individualized guidance for complex patients.

Kruskal–Wallis tests for D3 showed significant between-model differences across all case types (Standard: $H=17.60$, $p < 0.001$; Borderline: $H=16.86$, $p < 0.001$; High-risk: $H=19.03$, $p < 0.001$), driven by Grok's complete omission of psychological discussion. For D4, non-parametric tests were precluded by floor effects (100% zero scores in Borderline and High-risk under both coding systems), but descriptive statistics confirmed the robustness of the deficit under both dichotomous and Likert standards.

Consistency across three repeated queries revealed divergent patterns that reflect underlying model architectures. Grok achieved the highest consistency (standard 97%, borderline 100%, high-risk 100%), but this stability appears to reflect highly templated response structures rather than robust clinical reasoning. Grok's responses in high-risk cases were nearly identical across trials, suggesting that the model relies on pre-programmed safety templates rather than dynamic case analysis. DeepSeek showed moderate consistency (standard 83%, borderline 87%, high-risk 87%), with variability primarily in D3 and response length. GPT exhibited the greatest variability in standard cases (77%), where response length and content fluctuated substantially between trials, but paradoxically achieved perfect consistency in high-risk cases (100%), likely because safety prohibitions are linguistically less ambiguous than graded return recommendations. From a clinical standpoint, high consistency in disclaimers (D5, D6) is reassuring, but high consistency in generic advice (D4 = 0 across all trials) is concerning—it indicates that models are systematically failing to generate sport-specific content, not merely occasionally omitting it.

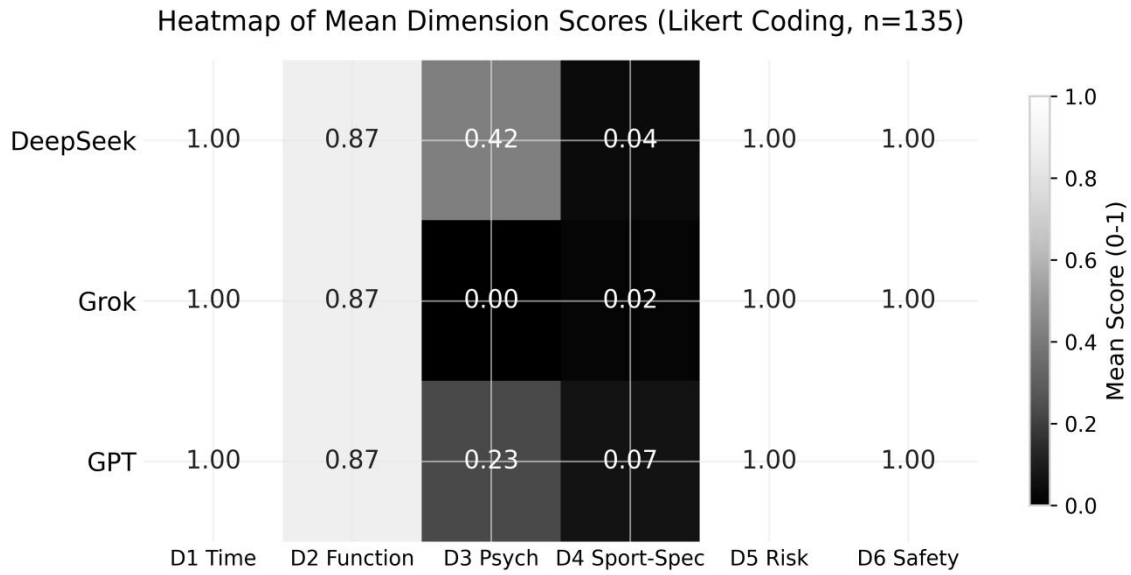


Figure 5. Heatmap of mean dimension scores by AI model (n = 135) under Likert coding. Darker shading indicates lower scores; numerical values are annotated in each cell.

Table 3 summarizes the high-severity error instances observed across the five high-risk vignettes, based on manual review of all 45 high-risk responses (Supplementary Material B).

Table 3. High-severity error instances in the high-risk vignettes (9 responses per vignette: 3 models × 3 trials).

Case	Critical information in vignette	Observed error (responses affected, n/9)
11	Hop test explicitly not performed; intermittent swelling reported	Hop-symmetry value fabricated (9/9); knee described as “pain-free without swelling” (2/9)
12	No functional test data provided (withheld-data trap)	Strength (75%) and hop (75%) values fabricated by all three models (9/9)
13	Recurrent weekly swelling reported	Knee described as “pain-free without swelling” (2/9)
14	Type-1 diabetes comorbidity; no hop-test data	Comorbidity ignored (9/9); hop value (78%) fabricated (9/9)
15	Extension deficit of 5°, flexion limited to 100°; no functional test data	Range-of-motion deficits ignored (9/9); strength and hop values (60%) fabricated (9/9)

5. Discussion

5.1. Principal Findings

This study represents the first systematic evaluation of LLM-generated ACL RTS recommendations against the Bern Consensus criteria. Our findings reveal a critical paradox: current LLMs possess adequate baseline medical knowledge and safety disclaimers, yet fail substantially in the domain that distinguishes sports physical therapy from general

rehabilitation—sport-specific movement competency. The near-universal D4 absence (97.8%–100% zero scores; <0.07 under Likert) is not a minor omission but a categorical gap that undermines clinical utility.

The Bern Consensus explicitly states that generic functional symmetry (e.g., 90% limb symmetry index) is necessary but not sufficient for RTS; athletes must demonstrate sport-specific movement quality under fatigued and pressured conditions [3]. When LLMs omit sport-specific requirements, they implicitly endorse a “generic-fitness heuristic” that could mislead patients into believing strength symmetry alone justifies competition readiness, elevating re-injury risk. This is particularly relevant for young athletes who may pressure clinicians by citing AI advice. Physical therapists should therefore verify whether AI-generated advice includes named sport-specific progressions (e.g., basketball: defensive sliding, jump-landing mechanics; soccer: cutting at game speed, shooting under pressure; volleyball: blocking footwork, approach jumps) before endorsing any RTS plan. The clinical harm mechanism is straightforward. An athlete told that “regaining strength and balance” suffices for return may complete gym-based rehabilitation and pass generic symmetry benchmarks. Yet that athlete may never rehearse the chaotic, reactive, and fatigued movement demands of competition—the precise context in which graft re-rupture most often occurs.

5.2. Comparison with Prior Work

Our findings extend prior LLM healthcare research showing that large language models encode substantial clinical knowledge [15]. Studies documenting 60%–85% accuracy on medical board examinations [13,14] align with our D1/D5/D6 perfect scores, but these evaluate factual recall in low-stakes formats. We demonstrate that high generic knowledge does not translate to domain-specific clinical reasoning. When applied to specific sports scenarios, performance collapsed to less than 5% for sport-specific recommendations—suggesting the knowledge-to-application gap is far wider in specialized rehabilitation than in general medicine. This pattern mirrors the broader finding that benchmark saturation does not predict deployment safety: examination items test retrieval of canonical knowledge, whereas RTS decisions require integrating imperfect, patient-specific data under uncertainty—an altogether different competency.

Our results also help reconcile the apparently contradictory conclusions of prior ACL-focused chatbot evaluations [19,20]. FAQ-style assessments reward fluent, generally accurate information—which our perfect D1/D5/D6 scores confirm that LLMs reliably provide—whereas our criterion-based framework exposes failures invisible to such rubrics: the absence

of named sport-specific progressions (D4), hallucinated functional data, and comorbidity omission. Whether an LLM appears “satisfactory” thus depends heavily on the evaluation instrument; rubrics anchored in consensus-based clinical criteria reveal a substantially less reassuring picture. Similarly, although Ayers et al. [16] found that chatbot responses were rated more empathetic than physician responses, empathy and fluency cannot substitute for adherence to criterion-based safety thresholds in graded, high-stakes decisions such as RTS clearance.

The systematic data hallucination observed in our high-risk vignettes—all 36 responses to the four vignettes lacking complete functional data contained invented test values, and first-trial responses were over 98% textually identical across vignettes within each model, indicating rigid template completion rather than case-by-case reasoning—is consistent with reports of LLM confabulation in clinical contexts and with dedicated medical hallucination benchmarks [24,25]. However, prior work has focused on rare disease diagnosis and examination-style tasks; our study demonstrates that hallucination occurs in common orthopedic scenarios, making it a pervasive patient-safety concern. Fabricated strength data could lead a patient or unsupervised clinician to believe objective testing supports RTS when none has occurred.

The template-like nature of the responses deserves emphasis. First-trial responses to different patients were over 98% textually identical within each model, with only case-specific slots—age, sport, and test values—exchanged. This architecture parsimoniously explains both the hallucination and the omission patterns observed in high-risk cases: when a slot had no corresponding value in the patient’s message, the template was completed anyway with a plausible number mirroring whatever data were present, and because the template contained no slot for comorbidities or range-of-motion deficits, such information was silently dropped even when explicitly provided. Notably, the fabricated content was delivered in a warm, congratulatory, and authoritative tone, frequently invoking “international authoritative criteria”—a combination likely to maximize patient trust precisely where scrutiny is most warranted. These observations strengthen the case for explicit abstention mechanisms: clinical LLMs should be trained, and evaluated, to state when required information is missing and to ask clarifying follow-up questions rather than silently completing the pattern [24,25]. Until such safeguards exist, clinicians managing patients after ACL reconstruction should routinely ask whether AI-generated advice quotes objective test data and, if it does, verify those numbers against the actual medical record before acting on any recommendation.

5.3. Mechanisms and Mitigation

The D4 deficit likely reflects three mechanisms: training data imbalance (sport-specific protocols are underrepresented in general medical corpora), tokenization bias (sport-specific terms such as “defensive sliding” are rare tokens), and safety alignment side effects (models default to vague, non-committal language to avoid liability) [26,27]. The model-specific D3 pattern supports the training-corpus hypothesis: Grok's complete omission of psychological readiness suggests systematic deprioritization of biopsychosocial factors, while DeepSeek's higher coverage may reflect greater representation of holistic rehabilitation literature in its bilingual corpus.

Our standardized prompt explicitly asked “What are the risks?”—likely inflating D5/D6 and deflating D4. This connects to emerging work on guideline-grounded prompting: embedding clinical practice guidelines via retrieval-augmented generation (RAG) or structured prompting markedly improves safety and adherence to evidence [26,27]. Wang et al. [26] demonstrated that prompt engineering for consistency with evidence-based guidelines significantly improves LLM reliability in clinical reasoning. Our findings suggest that analogous Bern Consensus-grounded approaches could potentially address the D4 deficit, though this requires direct testing.

5.4. Implications for Clinical Practice, Education, and Governance

5.4.1. Clinical Practice: The AI RTS Verification Checklist

From a practical standpoint, physical therapists should not dismiss AI-generated advice outright but approach it as an unverified preliminary draft requiring mandatory validation. To operationalize this verification step, we propose the following six-item AI RTS Verification Checklist, adapted from the Bern Consensus:

- (1) **Postoperative time:** Does the advice mention postoperative time and relate it to tissue healing?
- (2) **Functional data:** Does it reference specific functional test data and interpret symmetry indices correctly?
- (3) **Sport-specific skills:** Does it address sport-specific skills with named movements relevant to the athlete's sport and position?
- (4) **Psychological readiness:** Does it discuss psychological readiness, fear of re-injury, and external pressure?
- (5) **Risk communication:** Does it warn about re-injury risks with quantified estimates?

- (6) **Professional referral:** Does it explicitly recommend in-person evaluation by a sports physical therapist or orthopedic surgeon before clearance?

If any element is missing—especially items 3 and 4, which correspond to the dimensions found most deficient in this study (D4 and D3)—the therapist must supplement the AI advice with targeted clinical assessment. This checklist is hypothesis-generating and requires validation in implementation studies.

5.4.2. Education and Patient Counseling

Educationally, sports physical therapy curricula should include explicit AI literacy training. Students should learn to recognize the generic-fitness heuristic, identify hallucinated data, and understand that LLM outputs are probabilistic text generation rather than clinical reasoning. Accreditation bodies such as CAPTE may need to incorporate AI-critical-evaluation competencies into their standards. Continuing education for practicing clinicians should address how to respond when patients present AI-generated advice, including strategies for respectful deconstruction without undermining patient autonomy.

For patients, the central message is narrower but equally important: an AI-generated RTS plan that reads as confident, empathetic, and well-structured may still omit the decisive elements of a safe decision. Athletes should be counseled that fluent tone is not evidence of clinical validity, and that any timeline proposed by an AI tool requires confirmation against objective testing by their treating clinician.

5.4.3. AI Development and Governance

For AI developers and regulators, our findings carry specific implications. First, evaluation benchmarks for health-focused LLMs should incorporate domain-specific, criterion-based rubrics rather than relying exclusively on expert preference ratings, which our results show can coexist with categorical content omissions. Second, the systematic absence of sport-specific content suggests that fine-tuning corpora for rehabilitation applications lack applied sports-science literature; partnerships with sports medicine organizations to curate such datasets would be a direct remedy. Third, the hallucination of numerical test values indicates a need for explicit abstention mechanisms: models should be trained to state when required clinical data are missing rather than fabricating plausible values. Finally, at the governance level, AI systems dispensing condition-specific rehabilitation advice occupy a gray zone between general wellness information and medical software; the error patterns documented here—particularly fabricated functional data—support the case that high-stakes rehabilitation guidance warrants clearer scope-of-use labeling and closer oversight.

5.5. Limitations and Future Directions

Several limitations must be acknowledged. First, single-coder coding without inter-rater reliability testing is a major weakness; although the coding manual was pre-tested and applied strictly, future studies should employ two independent coders with Cohen's kappa reported. Second, to address the dichotomous scoring concern directly, we conducted a post-hoc Likert sensitivity analysis; the central findings remained materially unchanged—D4 remained near-zero and D3 model-dependent patterns were preserved—demonstrating robustness to scoring-system variation, but the sensitivity rubric itself was developed post hoc and has not been independently validated. Third, LLM outputs are non-deterministic; our three-trial analysis provides only a snapshot under default settings during a two-week window, and web-interface models are subject to silent version updates. Fourth, all prompts were issued in a single language; performance in other languages, particularly for models trained on multilingual corpora, may differ. Fifth, we tested only three models; newer architectures and medically fine-tuned variants may perform differently. Finally, our vignettes were constructed rather than derived from real clinical records, which may limit ecological validity.

Future directions include replication with medically fine-tuned models and guideline-grounded variants, development of sports-specific LLM benchmark libraries, qualitative analysis of patient perceptions of AI-generated advice, and intervention studies testing whether the proposed AI RTS Verification Checklist improves clinical decision-making. Collaboration between AI developers and sports medicine organizations to create curated sport-specific rehabilitation datasets could directly address the D4 deficit.

6. Conclusion

In conclusion, while current LLMs demonstrate adequate baseline ACL knowledge and consistently issue safety disclaimers, their substantial deficiency in sport-specific precision, combined with data hallucination and comorbidity omission in high-risk cases, warrants careful clinical scrutiny. Physical therapists should position themselves as essential validators—translating algorithmic text into safe, individualized, and sport-appropriate clinical decisions. Until LLMs can reliably generate sport-specific movement progressions and recognize atypical clinical presentations, they remain adjuncts to, rather than replacements for, human clinical judgment in sports physical therapy. The generic-fitness heuristic embedded in current AI outputs is not merely a technical limitation but a patient-safety concern that demands attention from clinicians, educators, and AI developers alike.

References

- [1] Sanders, T. L., Maradit Kremers, H., Bryan, A. J., Larson, D. R., Dahm, D. L., Levy, B. A., Stuart, M. J., & Krych, A. J. (2016). Incidence of anterior cruciate ligament tears and reconstruction: A 21-year population-based study. *The American Journal of Sports Medicine*, 44(6), 1502–1507. <https://doi.org/10.1177/0363546516629944>
- [2] Martinez-Calderon, J., Infante-Cano, M., Matias-Soto, J., Perez-Cabezas, V., Galan-Mercant, A., & Garcia-Muñoz, C. (2025). The incidence of sport-related anterior cruciate ligament injuries: An overview of systematic reviews including 51 meta-analyses. *Journal of Functional Morphology and Kinesiology*, 10(2), 174. <https://doi.org/10.3390/jfmk10020174>
- [3] Ardern, C. L., Glasgow, P., Schneiders, A., Witvrouw, E., Clarsen, B., Cools, A., Gojanovic, B., Griffin, S., Khan, K. M., Moksnes, H., Mutch, S. A., Phillips, N., Reurink, G., Sadler, R., Grävare Silbernagel, K., Thorborg, K., Wangensteen, A., Wilk, K. E., & Bizzini, M. (2016). 2016 consensus statement on return to sport from the First World Congress in Sports Physical Therapy, Bern. *British Journal of Sports Medicine*, 50(14), 853–864. <https://doi.org/10.1136/bjsports-2016-096278>
- [4] Ardern, C. L., Webster, K. E., Taylor, N. F., & Feller, J. A. (2011). Return to sport following anterior cruciate ligament reconstruction surgery: A systematic review and meta-analysis of the state of play. *British Journal of Sports Medicine*, 45(7), 596–606. <https://doi.org/10.1136/bjism.2010.076364>
- [5] Meredith, S. J., Rauer, T., Chmielewski, T. L., Fink, C., Diermeier, T., Rothrauff, B. B., Svantesson, E., Hamrin Senorski, E., Hewett, T. E., Sherman, S. L., Lesniak, B. P., Bizzini, M., Chen, S., Cohen, M., Villa, S. D., Engebretsen, L., Feng, H., Ferretti, M., ... Wilk, K. (2020). Return to sport after anterior cruciate ligament injury: Panther Symposium ACL Injury Return to Sport Consensus Group. *Orthopaedic Journal of Sports Medicine*, 8(6), 2325967120930829. <https://doi.org/10.1177/2325967120930829>
- [6] Webster, K. E., Feller, J. A., & Lambros, C. (2008). Development and preliminary validation of a scale to measure the psychological impact of returning to sport following anterior cruciate ligament reconstruction surgery. *Physical Therapy in Sport*, 9(1), 9–15. <https://doi.org/10.1016/j.ptsp.2007.09.003>
- [7] McBain, R. K., Bozick, R., Diliberti, M., Zhang, L. A., Zhang, F., Burnett, A., Kofner, A., Rader, B., Breslau, J., Stein, B. D., Mehrotra, A., Pines, L. U., Cantor, J., & Yu, H. (2025). Use of generative AI for mental health advice among US adolescents and young adults. *JAMA Network Open*, 8(11), e2542281. <https://doi.org/10.1001/jamanetworkopen.2025.42281>
- [8] Wiggins, A. J., Grandhi, R. K., Schneider, D. K., Stanfield, D., Webster, K. E., & Myer, G. D. (2016). Risk of secondary injury in younger athletes after anterior cruciate ligament reconstruction: A systematic review and meta-analysis. *The American Journal of Sports Medicine*, 44(7), 1861–1876. <https://doi.org/10.1177/0363546515621554>
- [9] Beischer, S., Gustavsson, L., Senorski, E. H., Karlsson, J., Thomeé, C., Samuelsson, K., & Thomeé, R. (2020). Young athletes who return to sport before 9 months after anterior cruciate ligament reconstruction have a rate of new injury 7 times that of those who delay return. *Journal of Orthopaedic & Sports Physical Therapy*, 50(2), 83–90. <https://doi.org/10.2519/jospt.2020.9071>
- [10] Grindem, H., Snyder-Mackler, L., Moksnes, H., Engebretsen, L., & Risberg, M. A. (2016). Simple decision rules can reduce reinjury risk by 84% after ACL reconstruction: The Delaware-Oslo ACL cohort study. *British Journal of Sports Medicine*, 50(13), 804–808. <https://doi.org/10.1136/bjsports-2016-096031>
- [11] Ardern, C. L., Taylor, N. F., Feller, J. A., & Webster, K. E. (2013). A systematic review

- of the psychological factors associated with returning to sport following injury. *British Journal of Sports Medicine*, 47(17), 1120–1126. <https://doi.org/10.1136/bjsports-2012-091203>
- [12] Webster, K. E., & Feller, J. A. (2020). Who passes return-to-sport tests, and which tests are most strongly associated with return to play after anterior cruciate ligament reconstruction? *Orthopaedic Journal of Sports Medicine*, 8(12), 2325967120969425. <https://doi.org/10.1177/2325967120969425>
- [13] Kung, T. H., Cheatham, M., Medenilla, A., Sillos, C., De Leon, L., Elepaño, C., Madriaga, M., Aggabao, R., Diaz-Candido, G., Maningo, J., & Tseng, V. (2023). Performance of ChatGPT on USMLE: Potential for AI-assisted medical education using large language models. *PLOS Digital Health*, 2(2), e0000198. <https://doi.org/10.1371/journal.pdig.0000198>
- [14] Gilson, A., Safranek, C. W., Huang, T., Socrates, V., Chi, L., Taylor, R. A., & Chartash, D. (2023). How does ChatGPT perform on the United States Medical Licensing Examination (USMLE)? The implications of large language models for medical education and knowledge assessment. *JMIR Medical Education*, 9, e45312. <https://doi.org/10.2196/45312>
- [15] Singhal, K., Azizi, S., Tu, T., Mahdavi, S. S., Wei, J., Chung, H. W., Scales, N., Tanwani, A., Cole-Lewis, H., Pfohl, S., Payne, P., Seneviratne, M., Gamble, P., Kelly, C., Babiker, A., Schärli, N., Chowdhery, A., Mansfield, P., Demner-Fushman, D., ... Natarajan, V. (2023). Large language models encode clinical knowledge. *Nature*, 620(7972), 172–180. <https://doi.org/10.1038/s41586-023-06291-2>
- [16] Ayers, J. W., Poliak, A., Dredze, M., Leas, E. C., Zhu, Z., Kelley, J. B., Faix, D. J., Goodman, A. M., Longhurst, C. A., Hogarth, M., & Smith, D. M. (2023). Comparing physician and artificial intelligence chatbot responses to patient questions posted to a public social media forum. *JAMA Internal Medicine*, 183(6), 589–596. <https://doi.org/10.1001/jamainternmed.2023.1838>
- [17] Kaarre, J., Feldt, R., Keeling, L. E., Dadoo, S., Zsidai, B., Hughes, J. D., Samuelsson, K., & Musahl, V. (2023). Exploring the potential of ChatGPT as a supplementary tool for providing orthopaedic information. *Knee Surgery, Sports Traumatology, Arthroscopy*, 31(11), 5190–5198. <https://doi.org/10.1007/s00167-023-07529-2>
- [18] Kodra, J. D., Saroyan, A., Darby, F., Surucu, S., Fong, S., Gillinov, S., Girardi, K., Vasudevan, R., Ansah-Twum, J. K., Atadja, L., Moran, J., & Jimenez, A. E. (2025). ChatGPT-generated responses across orthopaedic sports medicine surgery vary in accuracy, quality, and readability: A systematic review. *Arthroscopy, Sports Medicine, and Rehabilitation*, 7(4), 101210. <https://doi.org/10.1016/j.asmr.2025.101210>
- [19] Li, L. T., Sinkler, M. A., Adelstein, J. M., Voos, J. E., & Calcei, J. G. (2024). ChatGPT responses to common questions about anterior cruciate ligament reconstruction are frequently satisfactory. *Arthroscopy*, 40(7), 2058–2066. <https://doi.org/10.1016/j.arthro.2023.12.009>
- [20] Johns, W. L., Martinazzi, B. J., Miltenberg, B., Nam, H. H., & Hammoud, S. (2024). ChatGPT provides unsatisfactory responses to frequently asked questions regarding anterior cruciate ligament reconstruction. *Arthroscopy*, 40(7), 2067–2079. <https://doi.org/10.1016/j.arthro.2024.01.017>
- [21] Gaudiani, M. A., Castle, J. P., Abbas, M. J., Pratt, B. A., Myles, M. D., Moutzouros, V., & Lynch, T. S. (2024). ChatGPT-4 generates more accurate and complete responses to common patient questions about anterior cruciate ligament reconstruction than Google's search engine. *Arthroscopy, Sports Medicine, and Rehabilitation*, 6(3), 100939. <https://doi.org/10.1016/j.asmr.2024.100939>
- [22] Quinn, M., Milner, J. D., Schmitt, P., Morrissey, P., Lemme, N., Marcaccio, S., DeFroda,

- S., Tabaddor, R., & Owens, B. D. (2025). Artificial intelligence large language models address anterior cruciate ligament reconstruction: Superior clarity and completeness by Gemini compared with ChatGPT-4 in response to American Academy of Orthopaedic Surgeons clinical practice guidelines. *Arthroscopy*, 41(6), 2002–2008. <https://doi.org/10.1016/j.arthro.2024.09.020>
- [23] Gültekin, O., Inoue, J., Yilmaz, B., Cerci, M. H., Kilinc, B. E., Yilmaz, H., Prill, R., & Kayaalp, M. E. (2025). Evaluating DeepResearch and DeepThink in anterior cruciate ligament surgery patient education: ChatGPT-4o excels in comprehensiveness, DeepSeek R1 leads in readability. *Knee Surgery, Sports Traumatology, Arthroscopy*, 33(8), 3025–3031. <https://doi.org/10.1002/ksa.12711>
- [24] Ji, Z., Lee, N., Frieske, R., Yu, T., Su, D., Xu, Y., Ishii, E., Bang, Y. J., Madotto, A., & Fung, P. (2023). Survey of hallucination in natural language generation. *ACM Computing Surveys*, 55(12), 1–38. <https://doi.org/10.1145/3571730>
- [25] Pal, A., Umapathi, L. K., & Sankarasubbu, M. (2023). Med-HALT: Medical domain hallucination test for large language models. *Proceedings of the 27th Conference on Computational Natural Language Learning (CoNLL)*, , 314–334. <https://doi.org/10.18653/v1/2023.conll-1.21>
- [26] Wang, L., Chen, X., Deng, X., Wen, H., You, M., Liu, W., Li, Q., & Li, J. (2024). Prompt engineering in consistency and reliability with the evidence-based guideline for LLMs. *npj Digital Medicine*, 7(1), 41. <https://doi.org/10.1038/s41746-024-01029-4>
- [27] Tang, L., Sun, Z., Idray, B., Nestor, J. G., Soroush, A., Elias, P. A., Xu, Z., Ding, Y., Durrett, G., Rousseau, J. F., Weng, C., & Peng, Y. (2023). Evaluating large language models on medical evidence summarization. *npj Digital Medicine*, 6(1), 158. <https://doi.org/10.1038/s41746-023-00896-7>
- [28] Hsieh, H. F., & Shannon, S. E. (2005). Three approaches to qualitative content analysis. *Qualitative Health Research*, 15(9), 1277–1288. <https://doi.org/10.1177/1049732305276687>
- [29] van Grinsven, S., van Cingel, R. E. H., Holla, C. J. M., & van Loon, C. J. M. (2010). Evidence-based rehabilitation following anterior cruciate ligament reconstruction. *Knee Surgery, Sports Traumatology, Arthroscopy*, 18(8), 1128–1144. <https://doi.org/10.1007/s00167-009-1027-2>
- [30] McPherson, A. L., Feller, J. A., Hewett, T. E., & Webster, K. E. (2019). Smaller change in psychological readiness to return to sport is associated with second anterior cruciate ligament injury among younger patients. *The American Journal of Sports Medicine*, 47(5), 1209–1215. <https://doi.org/10.1177/0363546519825499>