

Open Book Test: Measuring OpenEvidence Performance on EdAI Certification Questions

Background: Recent research (Dec 2025) showed OpenEvidence achieves 100% accuracy on USMLE but only 31-41% on specialty board questions from the MedXpertQA dataset.

Methods:

- Generated 268 certification questions using EdAI from specific medical articles
- Removed 83 questions where OpenEvidence lacked access to source material
- Tested remaining 185 questions on OpenEvidence platform under "open book" conditions
- Documented accuracy across 7 medical specialties

Here's the specialty-by-specialty breakdown:

<i>Specialty</i>	<i>Correct</i>	<i>Total</i>	<i>Accuracy</i>
<i>Obesity Medicine</i>	19	24	79.17%
<i>Cosmetic Surgery (ABCS)</i>	24	32	75.00%
<i>Pediatric Neurosurgery</i>	30	40	75.00%
<i>–Gynecologic Oncology</i>	12	17	70.59%
<i>Urologic Oncology</i>	13	20	65.00%
<i>Laser Surgery (ABLS)</i>	18	32	56.25%
<i>Facial Cosmetic Surgery (ABFCS)</i>	11	20	55.00%
<i>Overall</i>	127	185	68.65%

Key Findings:

1. **Open book failure:** Despite direct access to source articles (83 questions excluded to ensure this), accuracy was only 68.65%
2. **Below passing threshold:** Medical students require 70%+ to pass certification

3. **Specialty variation:** 24-point spread (55-79%) demonstrates inconsistent reasoning
4. **Severe failures:** Two specialties (ABLS, ABFCS) failed nearly half of questions despite source access

Implications: General-purpose medical AI cannot serve professional certification markets. Content access ≠ content control. The ability to retrieve information does not equal the ability to reason through certification-level complexity.

EdAI Differentiation: Content-driven intelligence architecture generates questions that require synthesis and clinical reasoning beyond RAG retrieval capabilities. This creates a defensible moat in high-stakes certification domains.