



Standard Setting Report: Computed Tomography - Effective September 2026

Background

The mission of the American Registry of Radiologic Technologists (ARRT) is to “to promote safe, high-quality patient care through credentialing, collaboration, and advocacy.” Within ARRT’s examination program, standard setting is a formal, systematic process that links examination performance to a defensible determination of minimum competence for certification. The previous standard setting study for the Computed Tomography (CT) examination was conducted in March 2021.

This report details a standard setting study conducted in April 30 – May 1, 2026, for the CT examination. It describes (a) committee composition, (b) the definition of the minimally qualified candidate (MQC), (c) the standard setting methods and training provided, (d) the results from each method, and (e) the ARRT Board of Trustees’ decision and implementation plan.

The CT examination is designed to assess entry level competence. The passing standard is intended to support public protection by distinguishing candidates who have achieved the minimum required knowledge and skill from those who have not.

Panelists completed required confidentiality and conflict-of-interest attestations consistent with credentialing best practices for accreditation review.

Facilitators:

- Adisack Nhouyvanisvong, Ph.D., Senior Psychometrician at ARRT
- Sage Ro, Ph.D., Senior Psychometrician at ARRT

Committee Composition

ARRT selected subject matter experts (SMEs) from the volunteer database to maximize diversity in geography, work setting, and professional experience in Computed Tomography. When feasible, ARRT biases selection toward practitioners earlier in their career because the examination is designed to assess entry level competence. ARRT also invites the radiologist assigned by the American College of Radiology to the relevant examination committee.

For this study, twelve SMEs were invited. Of those, all twelve SMEs participated as panelists. The panel included eight female and four male members. Ten panelists reported hospital-based practice and two reported clinic-based practice. Six panelists were experienced technologists and six were entry level practitioners. Ten panelists were technologists, one was a radiologist, and one reported another role. Table 1 below summarizes demographic details.

Panel composition was reviewed to ensure appropriate representation across key practice characteristics, including experience level, work setting, and role.



Table 1. Committee Demographics

Member	Location	GENDER	INSTITUTION	EXPERIENCE	ROLE	CREDENTIALS
1	MA	M	Hospital	Experienced	Radiologist	M.D.
2	NH	F	Hospital	Entry Level	Technologist	R.T.(R)(CT)(ARRT)
3	NY	F	Hospital	Experienced	Technologist	R.T.(R)(MR)(CT)(ARRT)
4	WV	F	Hospital	Entry Level	Technologist	R.T.(R)(MR)(CT)(ARRT)
5	GA	F	Hospital	Entry Level	Technologist	R.T.(R)(T)(CT)(ARRT)
6	IN	M	Hospital	Entry Level	Technologist	R.T.(R)(CT)(ARRT)
7	CA	M	Hospital	Experienced	Technologist	R.T.(R)(MR)(CT)(ARRT)
8	AR	F	Clinic	Experienced	Technologist	R.T.(R)(M)(CT)(VI)(ARRT)
9	TN	F	Hospital	Entry Level	Technologist	R.T.(R)(CT)(ARRT)
10	MN	F	Hospital	Entry Level	Technologist	R.T.(R)(CT)(ARRT)
11	WY	F	Clinic	Experienced	Technologist	R.T.(R)(M)(CT)(BD)(ARRT)
12	PA	M	Hospital	Experienced	Other	R.T.(R)(CT)(ARRT)

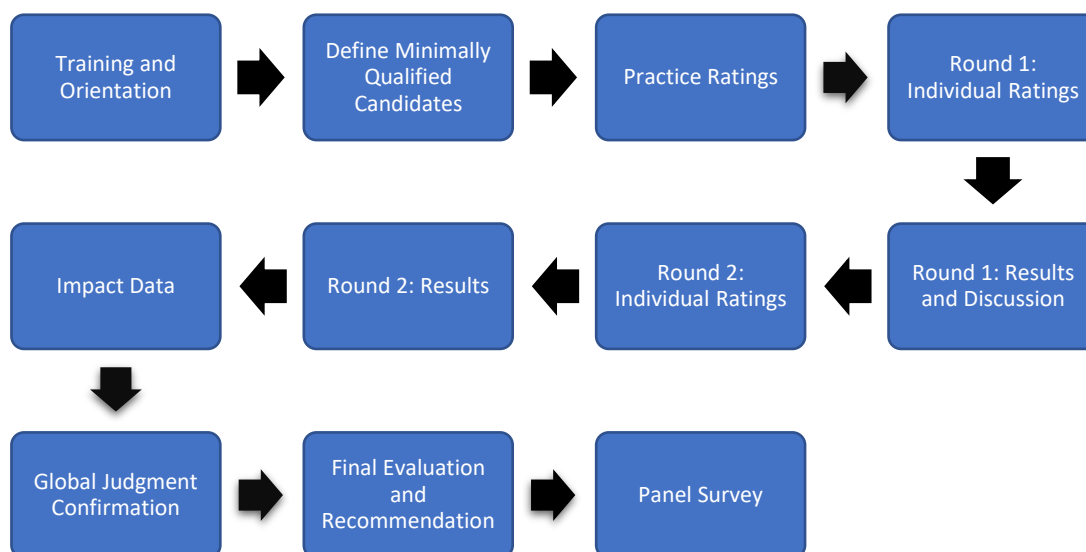
Standard Setting Process

Prior to the study, panelists were sent the examination content specifications, which describe the domain areas assessed on the CT examination. Panelists were also provided with task inventory, a list of activities performed by practicing CT technologists, and were asked to review both materials in advance.

Figure 1 summarizes the sequence of activities used during the study. The process began with an overview of standard setting and training on the Modified Angoff method (Angoff, 1971). Panelists then used the content specifications and task inventory to define the minimally qualified candidate, completed practice ratings, and conducted two rounds of Angoff ratings with structured discussion and feedback. Panelists were instructed to base their judgments on the minimally qualified candidates, rather than on clearly passing (highly experienced) candidates or failing candidates. Finally, panelists completed the global judgment confirmation activity (Beuk method) to evaluate the reasonableness of the Angoff-based results. The committee then made a final committee recommendation based on the Round 2 Modified Angoff results.



Figure 1. Procedures Followed during the Standard Setting Study



Minimally Qualified Candidate

Following training on the purpose of standard setting, its implications, and common sources of judgment biases, panelists developed a shared description of the minimally qualified candidate (MQC). The MQC is defined as a candidate who possesses the minimum knowledge, skills, and abilities required for safe and competent entry level practice in Computed Tomography. This candidate represents the lower boundary of acceptable performance, demonstrating just sufficient competence to practice safely and effectively.

Panelists reviewed the CT content specifications and task inventory and discussed what MQCs can perform reliably and independently, including the minimum level of clinical reasoning and procedural knowledge expected under typical testing conditions. Although the term “entry level” refers to the stage of professional practice, the MQC represents a performance threshold rather than a typical practitioner. Standard setting judgments are therefore based on the MQC as the point at which a candidate’s performance is just sufficient to meet the requirements for certification.

Methods

The study used a modified Angoff procedure (Angoff, 1971) as the primary item-centered standard setting method on a CT examination form. Global judgment confirmation method (Beuk, 1984) was then used to contextualize and evaluate the reasonableness of the Angoff-derived standard.

Modified Angoff Procedure

Prior to operational ratings, panelists completed a structured practice activity using a set of representative items. This exercise was designed to align their understanding of the MQC definition and the application of the rating scale, and to promote consistency in subsequent judgments. Panelists were instructed to make independent judgements, focusing on the MQC rather than on typical candidates or their own expectations.



In Round 1, panelists reviewed each item and estimated the percentage of MQCs who would answer the item correctly. No performance statistics were provided during Round 1.

Before Round 2, facilitators provided structured feedback to support informed reconsideration of judgments, including summary statistics of Round 1 ratings and identification of items with higher rating variability. Where available, first-time candidate proportion correct (or p-value) information was provided as interpretive context, with training that emphasized that operational data may inform (but should not replace) MQC-based judgment.

In Round 2, panelists reviewed their initial ratings and completed a second round of independent judgments. Final Angoff results were computed by summing item-level probabilities to produce a recommended raw score cut and associated percent correct cut on the standard setting examination form.

Following completion of Round 2 ratings, impact data were presented to the panel to provide context for decision-making. Specifically, the projected pass rate associated with the panel's recommended cut score was compared to the current operational pass rate based on recent candidate performance. This comparison allowed panelists to evaluate the practical implications of their judgments while maintaining focus on MQC expectations.

Global Judgment Confirmation (Beuk) Procedure

Following completion of the modified Angoff procedure, panelists provided independent global judgments. This procedure (Beuk, 1984) directed panelists away from item-level judgments and toward holistic evaluation of the examination form and expected candidate outcomes and is commonly used to contextualize Angoff results.

For the Beuk method, panelists were instructed to answer two questions designed to capture global expectations. Panelists indicated (a) the percentage of items an MQC would answer correctly to pass the examination and (b) the percentage of candidates who should be expected to pass. These judgments were compared to the Angoff results to evaluate consistency and support the reasonableness of the recommended passing standard. This step was confirmatory and did not serve to adjust the cut score.

Results

Results from the standard setting study are presented below, beginning with the modified Angoff procedure and followed by global judgment confirmation method. Tables 2 through 4 summarize the outcomes of each method and supporting evaluation information.

Table 2 presents the distribution of panelist Angoff judgments across rating rounds, including measures of central tendency, variability, inter-rater reliability, and standard error of judgment. Specifically, Table 2 shows that the Round 1 mean Angoff recommendation corresponded to a raw score of 110 (66.9%), with variability (8.0) across panelists. After feedback and discussion, Round 2 judgments converged to a higher mean raw score of 116 (70.2%), with reduced validity (7.9) and increased inter-rater reliability (0.91), indicating greater consistency among panelist's judgments.

Table 2. Modified Angoff Results



First Round			Second Round		
Statistic	Raw Scale	% Scale	Statistic	Raw Scale	% Scale
Mean:	110	66.9	Mean:	116	70.2
SD:	8.0	4.8	SD:	7.9	4.8
Min:	94.9	57.5	Min:	102.3	62.0
Max:	121.4	73.6	Max:	131.3	79.5
Sample SE:	2.3	1.4	Sample SE:	2.3	1.4
IRR-based SE:	4.2	2.5	IRR-based SE:	2.4	1.4
Max possible:	165	100	Max possible:	165	100
Inter-rater Reliability: 0.72			Inter-rater Reliability: 0.91		
Number of raters: 12			Number of raters: 12		

Note: SEJ denotes the standard error of judgment and reflects uncertainty (i.e., variability) in the panel's judgment-based standard. Inter-rater reliability (IRR) estimates describe agreement among panelists.

The projected impact of the recommended cut score (116 or 70.2% correct) was evaluated by comparing the associated pass rate (62% passing) to the current operational pass rate (73% passing). This comparison provided context for understanding potential changes in candidate outcomes and supported evaluation of the reasonableness and practical implications of the recommended standard.

Results from the Global Judgement Confirmation (Beuk method) are summarized in Table 3. Panelists' holistic judgments indicated a mean expected percent correct of 71.7% and a mean expected pass rate of 68.9%. These results were generally consistent with the Round 2 Modified Angoff outcomes and provided supporting context for evaluating the reasonableness of the recommended passing standard. Variability in panelists' global judgments, as reflected in the standard deviations, was within an acceptable range for this type of evaluation.

Table 3. Global Judgment Confirmation (Beuk) Results

Mean expected percent correct	71.7%
SD cut score	4.9%
Mean expected pass rate	68.9%
SD pass rate	6.3%

To evaluate the potential impact of judgment uncertainty on candidate outcomes, projected passing rates were examined across a range of cut scores defined by ± 1 and ± 2 standard errors of judgment (SEJ). Table 4 presents projected passing rates (62%) associated with the Round 2 Angoff cut score, as well as Global Judgment Confirmation (Beuk) scenarios (65%), relative to the current operational passing standard (73%).

As shown in Table 4, projected pass rates vary across methods and SEJ-adjusted cut scores, illustrating the sensitivity of candidate outcomes to modest changes in the passing standard. The table also provides a reference comparison to the current operational cut score, only supporting evaluation of continuity and potential impact.



Table 4. Panelist Cut Scores and Projected Passing Rate

Projected Cut and Passing Rate	Cut -2SEJ	Cut - 1SEJ	Cut	Cut + 1SEJ	Cut + 2SEJ
Angoff Cut:	112	114	116	119	121
Pass rate:	69%	65%	62%	57%	52%
Global Judgment (Beuk)	110	112	114	117	119
Pass rate:	73%	65%	65%	60%	57%
Current Cut			110		
Pass rate:			73%		

Note: Global judgment values are presented for comparison purposes only and do not represent an alternative passing standard.

The current cut score shown in Table 1 reflects the raw score equivalent to the standard setting form required to maintain the currently approved passing standard. Although the Board approved a passing standard corresponding to 111 items correct in 2021, application of that same standard to the form used in this standard setting study results in a raw score equivalent near 110 items correct. In other words, while the underlying passing standard remains the same, the raw score required to meet that standard may change across examination forms because forms differ in difficulty and statistical characteristics.

Board of Trustees' Decision and Implementation

The ARRT Board of Trustees reviewed the results and discussed the impact of potential new standards before approving a final standard for the Computed Tomography exam.

After review and discussion, the Board approved a new standard equivalent to 112 out of 165 items on the exam form used for this meeting. The new standard will go into effect September 2026 and remain in place until the next standard setting is scheduled to take place. ARRT staff expect a future pass rate for first-time candidates around 69% based on the impact data provided to both the Board and standard setting committee.

References

Angoff, W. H. (1971). Scales, norms, and equivalent scores. In R. L. Thorndike (Ed.), *Educational measurement* (2nd ed, pp. 508-600). Washington, DC: American Council on Education.

Beuk, C. H. (1984). A method for reaching a compromise between absolute and relative standards in examinations. *Journal of Educational Measurement*, 21(2), 147-152.

