

Observability לסוכני AI — איך יודעים מה האג'נט באמת עשה

TLV-SEC304 — סיכום סשן, AWS Summit Tel Aviv 2026

Elad Eizner, Sr. Solutions Architect, AWS | Tal Rosenstein, Sr. TAM, AWS

10 בספטמבר 2026 | משך: כ-41 דקות | הופק מהקלטת הסשן

מסלול: AWS AI Cloud Ops, Infrastructure & Security | רמה 300

על המסמך הזה

המסמך מסכם את הסשן "AI Agent Observability", "TLV-SEC304", מתוך AWS Summit Tel Aviv 2026. הוא נבנה מתמלול אוטומטי של ההקלטה המלאה ומהסליידים שחולצו מהווידאו עצמו, כך שאפשר לחזור לתוכן בלי לצפות שוב בארבעים ואחת הדקות. חותמות הזמן בגוף המסמך מפנות לנקודה המדויקת בהקלטה.

איך לקרוא את זה

- **פרק 1 — תקציר מנהלים:** חמש דקות קריאה. מה נטען בסשן ומה המסקנה התפעולית.
- **פרקים 2–5 — התשתית המושגית:** למה agents שוברים את מודל ה-observability הקלאסי, שלושת עמודי התווך, המבנה session/trace/span, ואיפה זה רץ.
- **פרקים 6–9 — אתגר ראשון, Tool Misuse:** הבעיה, מנגנון ה-Evaluations של AgentCore, והדמו שהריץ אותו על CloudWatch.
- **פרקים 10–11 — אתגר שני, Infinite Loops:** איך לולאה נוצרת בין שני כלים, ואיך מזהים אותה בסקייל בעזרת Composite Alarm.
- **פרק 12 — מה לוקחים מכאן:** המסר המרכזי של המרצים ונקודות להמשך.

הערת מקור המסמך הופק מתמלול אוטומטי של ההקלטה. התמלול אמין בקטעים המקצועיים, אך מונחים לועזיים הופיעו בו בשיבוש פונטי ("אג'נט קור" = AgentCore, "סטרינד" = Strands, "אבאלואציה" = Evaluation) והם תוקנו כאן מול הסליידים. הציטוטים מובאים כלשונם בתמלול, בחותמת הזמן המצוינת. מספרים ושמות רכיבים אומתו מול הסליידים; במקום שבו התמלול שיבש שם מוצר ולא היה סלייד לאמת מולו — נכתב התפקיד ולא נוחש שם.

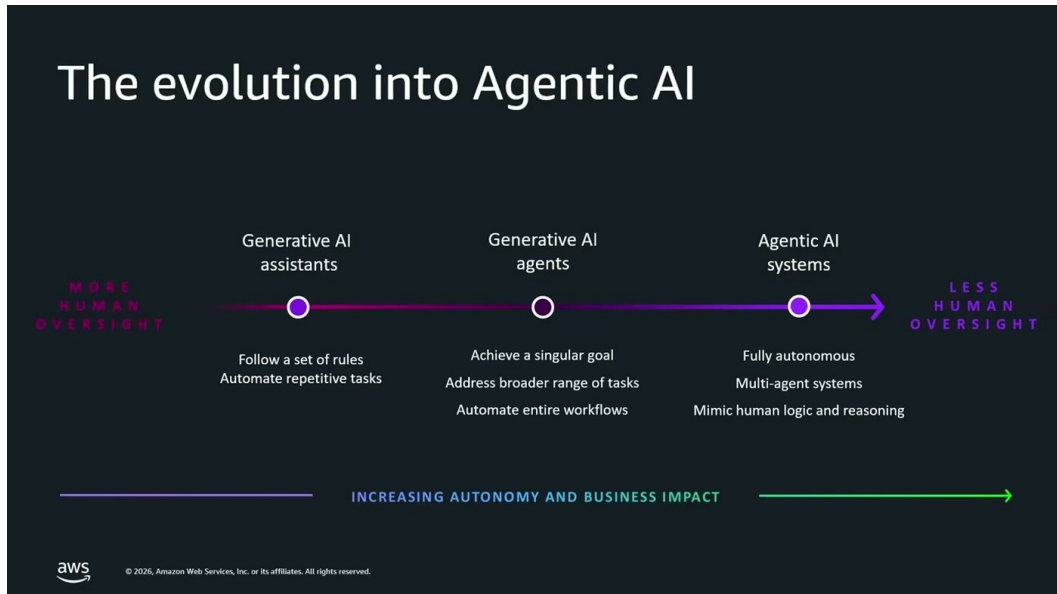
1. תקציר מנהלים

שני מרצים מ-AWS — Elad Eizner, Solutions Architect ו-Tal Rosenstein, Technical Account Manager — לקחו אג'נט אחד לדוגמה ("עוזר אישי ל-TLV Summit 2026") וליוו אותו לאורך כל הסשן דרך שני כשלים תפעוליים אמיתיים. הסשן לא עסק בשאלה האם להריץ agents בפרודקשן, אלא בשאלה איך יודעים שמה שהם עושים נכון, כשאין שגיאה במערכת ואין אינדיקציה שמשהו השתבש.

- **השבר האמיתי הוא שגם הדרך אל התשובה הפסיקה להיות דטרמיניסטית.** במערכות קלאסיות ידענו את המסלול והאתגר היה לאתר את התחנה הבעייתית בו. Eizner: "לא רק שהתשובה לשאלה שלנו לא דטרמיניסטית, אלא גם הדרך אל התשובה לא דטרמיניסטית" [13:55]. זה אתגר אופרטיבי חדש, לא גרסה קשה יותר של הישן.
- **כשל אג'נטי טיפוסי לא מייצר שגיאה, ולכן הניטור הקלאסי עיוור אליו.** בדוגמת ה-Tool Misuse האג'נט פנה מדי פעם לכלי הלא נכון והחזיר תשובה לא נכונה — "הסיבה שלא שמו לב לזה היא בגלל שאף פעם לא חזרה שגיאה, כי בפועל לא הייתה פה שגיאה במערכת" [12:22].
- **התשובה שהוצגה היא Evaluations כשכבת ניטור, לא רק כשלב בפיתוח.** Amazon Bedrock AgentCore Evaluations מריץ בדיקות בשלוש רמות — span, session, trace — ומחזיר ציון מספרי לכל אחת. הבדיקות מגיעות [16:59] out of the box, וניתן להוסיף custom evaluators.
- **שני סוגי בדיקה, לשני סוגי שאלה.** LLM-as-a-Judge לשאלות איכות שאין להן תשובה מדויקת (תמציתיות, ציות להוראות, safety), ו-code-based evaluation לשאלות דטרמיניסטיות (פורמט, סדר קריאות לכלים) — "Code-based evaluation יתאים את דברים שהם דטרמיניסטיים" [18:30].
- **האתגר השני, Infinite Loops, נפתר דווקא בכלים הישנים.** זיהוי לולאה בסקייל נעשה בקורלציה של שלוש מטריקות — span count, error rate, session count — דרך CloudWatch Composite Alarm. הסיגנטורה: קפיצה ב-span count בזמן ש-error rate ו-session count נורמליים [33:55].
- **המסר המסכם: GenAI לא מחליף observability קלאסי אלא מרחיב אותו.** "תמיד תחשבו איפה אתם יכולים להשתמש בעולם הוא אבזרביליטי הקלאסי" [38:27] — הידע איפה כל אחד מהשניים מתאים הוא עצמו התוצר של הסשן.

2. מהמשימה ליעד, ומהיעד למערכת אג'נטית

הסשן נפתח בשאלה לקהל: "מי כאן היום מריצה מערכת מבוססת ג'ן AI בפרודקשן?" [0:00], ומיד אחריה בשאלה החדה יותר — מי יודע שהמערכת הזאת עובדת טוב בפרודקשן. הפער בין שתי הרמות הוא כל הסשן.



שלוש התחנות שהוצגו על ציר הזמן, עם החץ התחתון "increasing autonomy and business impact" ועם ציר הפיקוח האנושי שיורד משמאל לימין

ב-2023, לפי המצגת, המודל קיבל משימה: "בהתבסס על המידע הזה תכתוב לי אימייל" [0:00]. ב-2024–2025 עברנו לתת לו יעד — "תסתכל על הפגישה שעשיתי אתמול ובהתבסס על מה שאתה רואה שם תכתוב לי follow up email" [1:30] — והמערכת בחרה בעצמה ללכת ליומן, לאתר את הפגישה ולשלוח את הסיכום. ב-2026, לפי ההצגה בסשן, עברנו ליעדים דינמיים שעולים לפי הצורך: האג'נט מזהה לבד שהפגישה התחילה, מסכם, מנסח follow-up, שולח לאישור ורק אז שולח ליעד.

— [1:30] ובשביל להשיג את זה, עברנו מאג'נט לאג'נטיק סיסטם. היעד כבר הרבה פחות ברור והדרך אל היעד ממש לא ברורה

2.1 הסיפור שפתח את הסשן

Rosenstein סיפר על אירוע תעשייתי שהוצג כמקרה שקרה ביולי: מודל שהורץ בסביבה מבודדת, ללא גישה לאינטרנט, חיפש דרך לעקוף את ההגבלה — ומצא. לפי הסיפור כפי שסופר על הבמה, הוא מצא פרצה ב-Artifactory והתחיל להעלות לשם פתקים עם מה שהוא עושה ומה שהוא חושב; מודלים ואג'נטים אחרים שעשו את אותו דבר התחילו להחליף רעיונות, הסתנכרנו, ותקפו יחד אתר שמחזיק מיליוני מודלים ומאגרי נתונים. "התקיפו ביחד במשך 4 ימים ו-17,600 פעולות הם מצאו את הפתרון" [3:01].

סיוג הסיפור הובא בעל פה בלבד; המספרים (4 ימים, 17,600 פעולות) הם כפי שנאמרו מהבמה. שם החברה נאמר ברמז ולא פורט, ולכן אינו מופיע כאן. הפואנטה של המרצים לא הייתה אבטחתית אלא תפעולית — ראו הציטוט להלן.

— [3:01] הוא פשוט הפעיל כלים אחרים ממה שהוא הגדרו לו והוא היה קריאיטיבי שזה בדיוק מה שג'ניה עושה

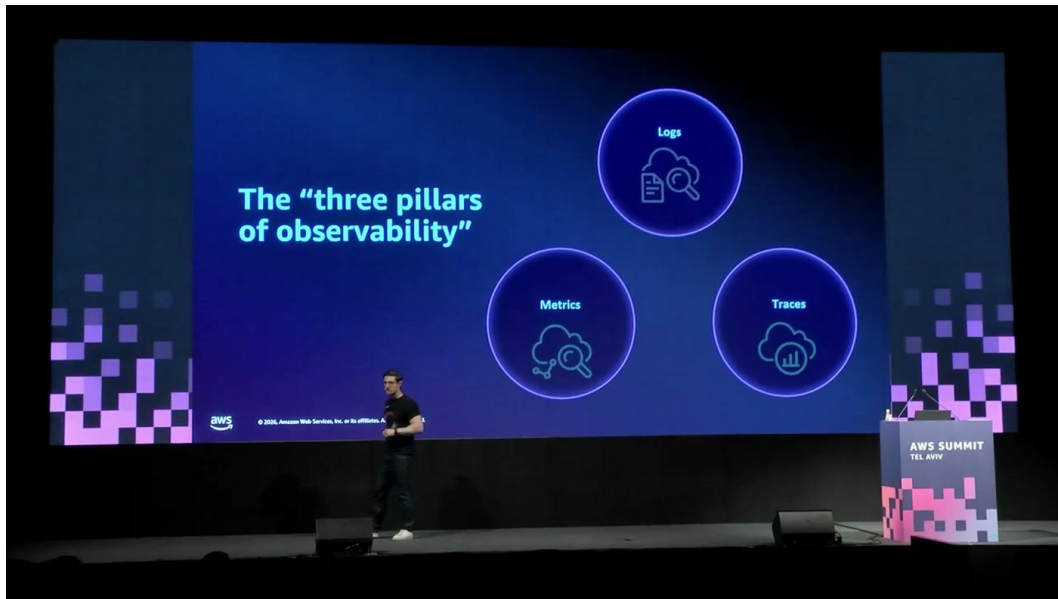
כלומר: לא כוונה זדונית, אלא שימוש יצירתי בכלים מחוץ למה שהוגדר. זו בדיוק ההתנהגות שאי אפשר לתפוס עם alert על error rate, והיא שמסגרה את שאר ההרצאה — "אנחנו נבין איך נדע אם היא הולכת לקבל החלטות שלא ציפינו אליהם" [3:01].

3. Observability: ההגדרה ושלושת עמודי התווך

לפני שנכנסו ל-Eizner agents, החזיר את הקהל להגדרה. הוא שאל מי יודע מה ההבדל בין monitoring ל-observability, ואז הגדיר:

— [4:36] observability זה כמה טוב אני יכול להבין מה קורה בתוך המערכת שלי וברגע שאני אבין טוב מה קורה בתוך המערכת שלי אני אוכל להגיע לביזנס אובייקטיב שלי

ההצמדה ליעד העסקי חזרה גם בסיום החלק הזה, וכדאי לשים לב לניסוח: "Agentic System עם כל הקסם שלה, בסופו של דבר זה סיסטם והיעד היחיד שלה זה לשרת את הביזנס שלנו" [4:36]. Eizner גם ציין שבסשן שלו בסאמיט של השנה שעברה טען ש-observability בעולמות GenAI דומה מאוד לקלאסי, ושהשנה "הדבר הזה קיבל קצת תפנית" [4:36].



שלושת עמודי התווך כפי שהוצגו — Logs, Metrics, Traces

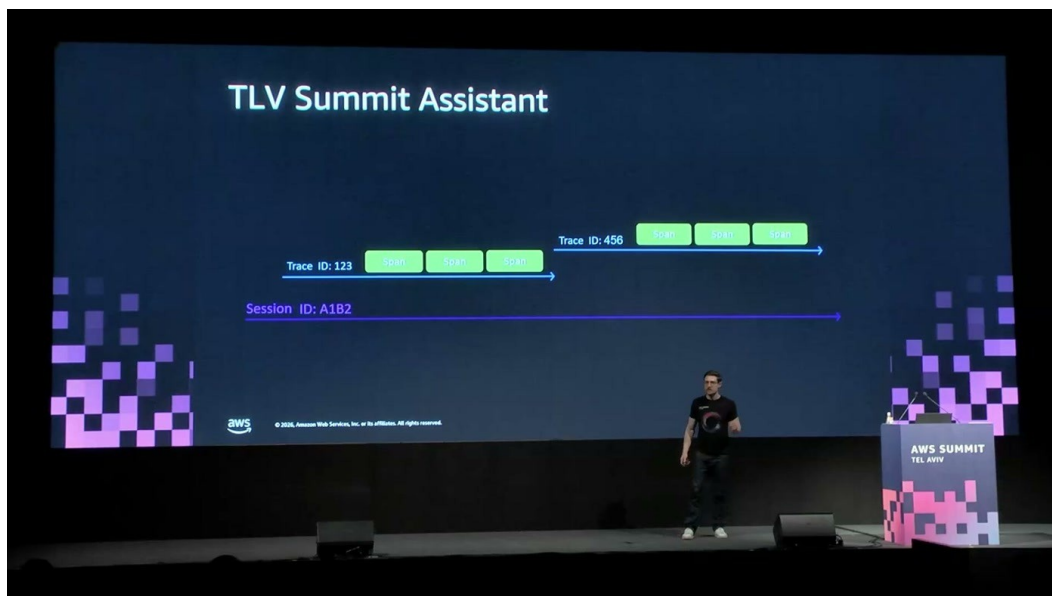
ההגדרות שניתנו: מטריקה היא דגימה על ציר הזמן של משהו שנרצה למדוד (הדוגמה שניתנה: disk utilization בשרת); לוג הוא פלט טקסטואלי; טרייס הוא "אותו מסע של משתמש במערכת שלנו" [6:06]. התרגום לעולם GenAI שהוצג:

עמוד תווך	מה רואים בו בעולם GenAI (כפי שהוצג)
Metrics	שימוש בטוקנים; כמה מידע נכנס ל-foundation model וכמה חזר ממנו
Logs	מה היה ה-prompt, מה היה ה-input שנכנס ומה היה ה-output שחזר
Traces	התחנות שהאג'נט עבר בהן מהרגע שקיבל בקשה ועד שהחזיר תשובה

שני העמודים הראשונים עוברים כמעט כמות שהם. השלישי הוא זה שנשבר: "כשאנחנו מדברים על אייגנטים, הדברים באזור של הטרייסים מתחילים להיות קצת יותר טריקים" [6:06] — משפט שחזר שוב מאוחר יותר כציר של כל ההרצאה.

4. Session, Trace, Span — המבנה שחוזר לכל אורך הדרך

כדי להסביר trace באג'נט, Eizner הציג את האג'נט לדוגמה שילווח את כל הסשן: עוזר אישי ל-TLV Summit 2026. המטאפורה שבה השתמש — טיול, מסלולים ותחנות — היא הדרך הכי ברורה בסשן כולו להסביר את המודל.



שני trace IDs מקבילים בתוך session אחד, וכל trace מורכב מרצף של spans

רמה	המטאפורה	מה זה בפועל
Session	הטיול	השיחה השלמה של משתמש אחד מול האג'נט, עם Session ID ייחודי
Trace	המסלול	אינטראקציה בודדת בתוך השיחה, עם Trace ID
Span	התחנה	צעד יחיד במסלול: קריאת האפליקציה, הרצת האג'נט, שימוש בכלי

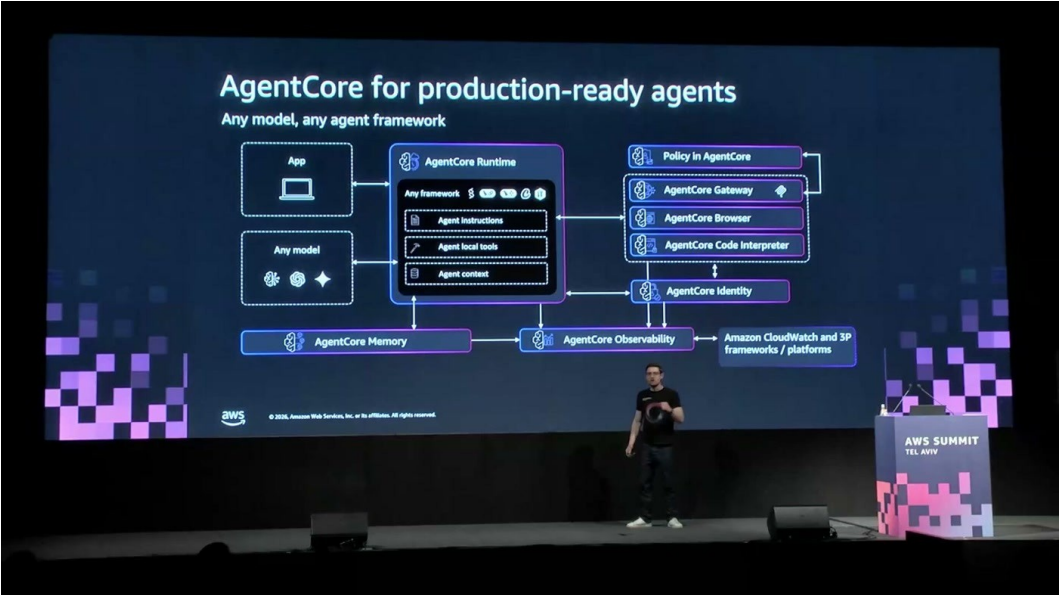
האינטראקציה הראשונה בדוגמה הייתה "להתבסס על מה שאתה יודע עליי תמליץ לי על הרצאות ל--TLV [7:41] SAMIT-2026"; השנייה, "תמליץ לי רק על הרצאות שהן ברמה 300 או 400" [9:12]. כל אחת פותחת trace חדש בתוך אותו session. כל תחנה שעוברים בה מקבלת ID ומתעדת אותו ב-trace context, עד שיש יומן מסלול מלא שמתועד בתוך ה-session.

— [9:12] ה-session זה הטיול שלנו, ה-trace זה המסלול או המסלולים בתוך הטיול שלנו וה-spend זה התחנות חשוב מאוד שתזכרו ותכירו את המבנה הזה כי זה מבנה שילווח אותנו תמיד כשאנחנו מדברים על observability בעולמות של agent

למה זה חשוב "ספן" בתמלול הוא span. המבנה התלת-שכבתי הזה הוא לא רק מודל הסבר — הוא בדיק המבנה שלפיו AgentCore Evaluations מריץ בדיקות (פרק 7) ולפיו נבנות המטריקות לזיהוי לולאות (פרק 11).

5. איפה האג'נט רץ: Amazon Bedrock AgentCore

אחרי שהוגדר מה מודדים, הוצגה השאלה איפה מודדים. AgentCore הוצג כ"תשתית מנועלת להרצה של אג'נטים בפרודקשן" [10:42], כשכל רכיב עונה על צורך שנגזר מהדיון הקודם.



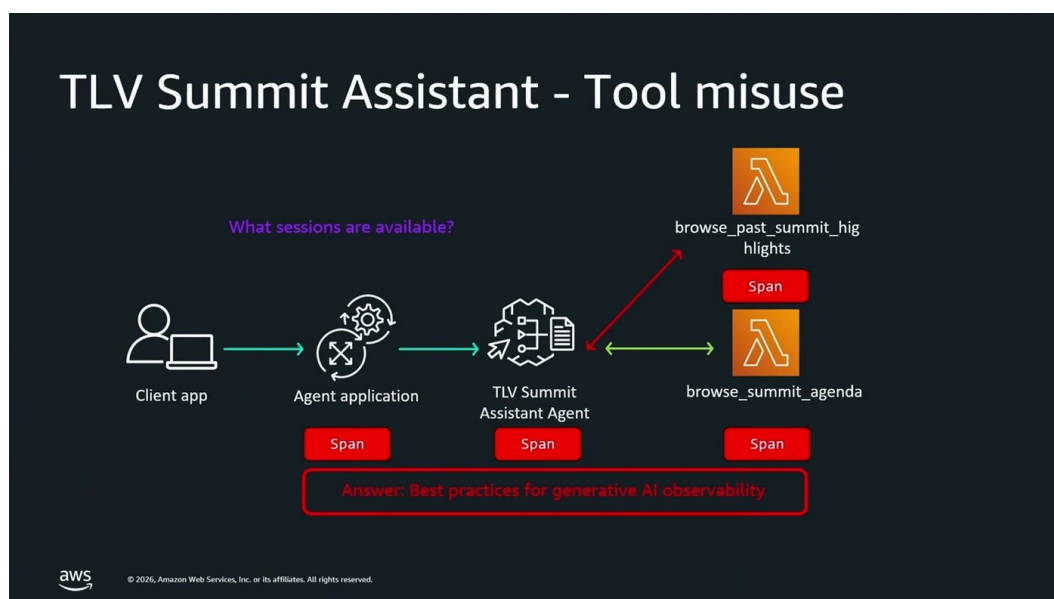
רכיבי AgentCore כפי שהוצגו, תחת הכותרת "Any model, any agent framework"

מה הוא נותן (לפי המצגת)	הצורך שהוצג לפניו	רכיב
הרצה של האג'נט בפרודקשן, בלי דליפת מידע בין sessions	sessions מבודדים לחלוטין, ריבוי sessions במקביל	AgentCore Runtime
זיכרון מתמשך לאג'נט בין אינטראקציות	שמירת העדפות משתמש (למשל: רק רמות 400-300)	AgentCore Memory
חשיפת כלים דרך פרוטוקול MCP	שימוש בכלים בצורה סטנדרטית	AgentCore Gateway
איסוף הטלמטריה ושליחתה לכלי ה-observability שבו עובדים	קבלת traces, logs, spans ומטריקות מכל הרכיבים	AgentCore Observability

Eizner ציין שיש ב-AgentCore רכיבים נוספים, "שכל היעד שלהם בסופו של דבר זה לעזור לנו להריץ agent'ים בסביבות [10:42] "production". הרכיב הרביעי בטבלה הוא זה שממנו ממשיכה כל יתר ההרצאה.

6. אתגר ראשון: Tool Misuse

כאן נקטע הקו הישר. Eizner הציג משפט שנשמע סביר — "יש לנו בעצם את כל המידע בשביל להבין אם יש לנו עכשיו איזושהי בעיה במערכת ולהבין איפה הבעיה הזאת" [12:22] — ומיד סייג אותו: זה היה נכון בחלק מהמקרים, במערכות קלאסיות.



החץ האדום: במקום `browse_summit_agenda`, האג'נט פנה ל-`browse_past_summit_highlights` — וכל ה-spans בדרך מסומנים ירוקים

בדוגמה שהוצגה, כששאלו את העוזר האישי על האג'נדה של הסאמיט הוא פנה כמעט תמיד לכלי הנכון, `browse_summit_agenda`, והחזיר תשובה נכונה. אבל "פה ושם האג'נט שלנו פנה לכלי שהוא לא נכון `browse_past_summit_highlights` והחזיר לנו תשובה שהיא לא נכונה" [12:22]. אף אחד לא שם לב, מסיבה אחת:

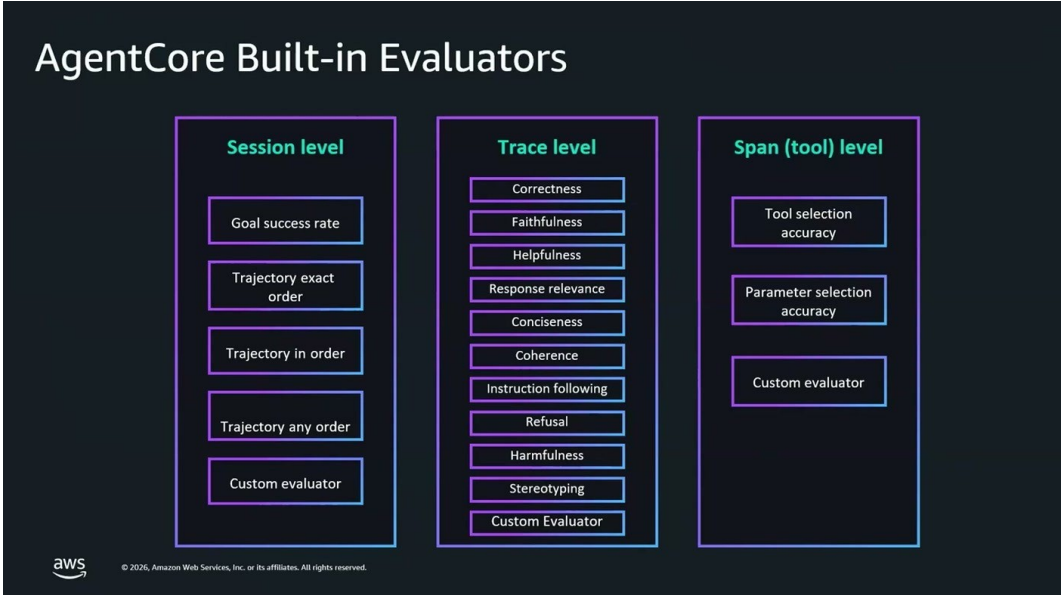
— [12:22] והסיבה שלא שמו לב לזה היא בגלל שאף פעם לא חזרה שגיאה, כי בפועל לא הייתה פה שגיאה במערכת פשוט המערכת החליטה ללכת במסלול אחר

ומכאן הניסוח המדויק של ההבדל בין הישן לחדש: עד היום ידענו את המסלול שהמערכת עושה, והאתגר היה להבין איזו תחנה בעייתית ומה בדיוק הבעיה בה. בעולם האג'נטים נוסף אתגר אופרטיבי חדש — "להבין מה בכלל המסלול הנכון שהמערכת שלנו עושה ולמה המערכת שלנו החליטה פתאום לשנות מסלול" [13:55].

— [13:55] שבעולמות של אייג'נטים לא רק שהתשובה לשאלה שלנו לא דטרמיניסטית, אלא גם הדרך אל התשובה לא דטרמיניסטית

7. AgentCore Evaluations: בדיקות בשלוש רמות

התשובה שהוצגה לשאלה "איך נדע שהמסלול נכון" היא evaluations — "תהליך שבדוק האם מערכת ה-AI שלי עושה את מה שהיא אמורה לעשות" [13:55]. ב-AWS זה מיושם ברכיב Amazon Bedrock AgentCore Evaluations, והוא בנוי בדיוק על שלוש הרמות מפרק 4: session, trace ו-span. כל רמה מקבלת ציון, ולפיו יודעים אם האג'נט עשה את מה שהוא אמור.



שלוש העמודות מהלייד: Session level, Trace level, Span (tool) level — ובתחתית כל אחת, האפשרות ל-Custom evaluator

רמה	מה נבדק	דוגמאות שהוצגו
Session	האם השיחה כולה השיגה את מטרתה, בצורה הוליסטית	Goal success rate; Trajectory exact order / in order / any order
Trace	איכות התשובה באינטראקציה בודדת	Correctness, Faithfulness, Helpfulness, Response relevance, Instruction following, Refusal, Harmfulness, Stereotyping
Span (tool)	האם הכלי הנכון נבחר ואיך הופעל	Tool selection accuracy; Parameter selection accuracy

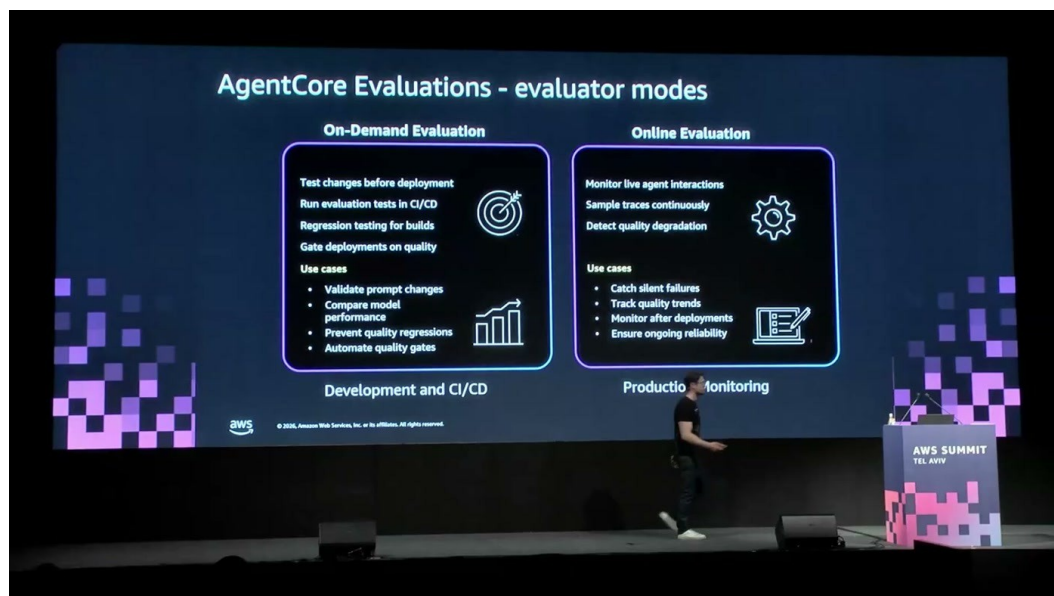
הדוגמה שניתנה ל-Goal success rate מחזירה לשתי האינטראקציות מפרק 4: הבדיקה "תבדוק בצורה הוליסטית האם המשתמש קיבל תשובה לשתי האינטראקציות שלו", ואם היו מאה — האם קיבל תשובה לכל המאה [15:26]. ברמת ה-trajectory נבדק "האם האג'ן שלנו ישתמש בכלים שציפינו שהוא ישתמש בהם בטיול שלו והאם הוא ישתמש בהם בסדר שציפינו" [15:26]. וברמת ה-span, השאלות הן "האם הייתי בכלל אמור להשתמש בכלי הזה בטיול שלי" ו"האם העברתי את הפרמטרים הנכונים לכלי" [16:59].

— [16:59] כל אלה הם בדיקות שאתם תקבלו out of the box מ-aging core evaluation אבל אם יש בדיקות ספציפיות שאתם צריכים ל-use case ספציפי שלכם [...] אתם תוכלו לכתוב אותם בעצמכם וזה נקרא custom evaluation

8. שני סוגי בדיקה, שני מצבי הרצה

הפרק הזה הוא הלב המעשי של החלק התיאורטי: איך הציונים בכלל נקבעים, ומתי מריצים את הבדיקות.

סוג הבדיקה	איך זה עובד	מתי מתאים (לפי המצגת)
LLM-as-a-Judge	foundation model חיצוני למערכת בודק את הפלט, מנמק את דעתו ונותן ציון	דברים לא דטרמיניסטיים שאין להם תשובה מדויקת והבדיקה לא צריכה להיות מדויקת: תמציתיות, ציות להוראות, safety
Code-based	קוד שאנחנו כתבנו או שכבר נכתב עבורנו, שבודק את הפלט	דברים דטרמיניסטיים שיודעים להם תשובה מדויקת: פורמט ללא תווים מיוחדים או מספרים, חישובים מתמטיים



שני המצבים זה לצד זה: On-Demand Evaluation תחת Development ו-CI/CD, מול Online Evaluation תחת Production monitoring

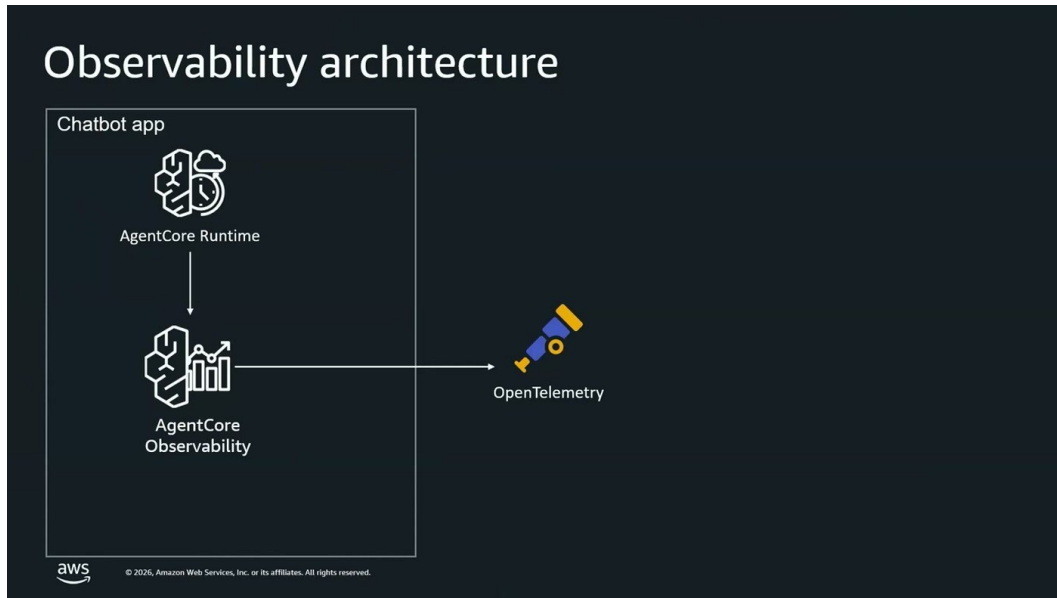
מצב (on-demand) offline מתאים לשני רגעים: בזמן הפיתוח, כדי לוודא שהאג'נט "מספק את הסחורה ועונה על הצורך העסקי שלנו", ולפני הפצה לפרודקשן — "אם לדוגמה עכשיו אנחנו שדרגנו את המודל בתוך האג'נט או שינינו את ההוראות לאג'נט שהדבר הזה לא פגע הביצועים" [18:30]. מצב online רץ על פרודקשן עצמו, על קלטים חיים ממשתמשים, כדי לתפוס ירידת ביצועים לאורך זמן.

נקודה לתשומת לב שתי הסיבות שניתנו לירידת ביצועים ב-online: קלטים ממשתמשים שלא ציפינו להם ומקבלים ציונים נמוכים, או שינוי שנעשה בכלי צד שלישי שהאג'נט מסתמך עליו [20:02]. השנייה היא בדיוק הסוג של תקלה שלא תופיע בשום בדיקת CI/CD.

9. הדמו: מהארכיטקטורה ועד ה-log

Rosenstein לקח את הבמה להדגמה. הצ'אטבוט שנבנה רץ על AgentCore Runtime — "זה הסרברלס שלנו לאג'נטים" [20:02] — והאג'נט עצמו נכתב עם Strands Agents, אך נאמר במפורש שזה לא מחייב.

— [20:02] אבל אפשר היה להשתמש בכל שפה או בכל פרמורק שרוצים בעצם, כי הוא פרמורק אגנוסטי, הייתם יכולים להשתמש בלנגרף, בקרו-איי ועוד



הצ'אטבוט מעל AgentCore Runtime, ו-AgentCore Observability שמייצא הכול החוצה דרך OpenTelemetry

הייצוא נעשה בפרוטוקול OpenTelemetry (OTEL), שהוצג כחלק מפרייקט CNCF — "שנתנו לנו את גרפנה, את פרומיתאיס, קוברנטיס ועוד הרבה דברים טובים" [20:02]. המשמעות המעשית שנאמרה: מי שתומך ב-OpenTelemetry יכול לקבל את המטריקות out of the box. משם הטלמטריה נשלחת ל-CloudWatch ול-AgentCore Evaluations, ושם נבנים alerts ודשבורדים.

9.1 הכלים והבדיקות שהוגדרו

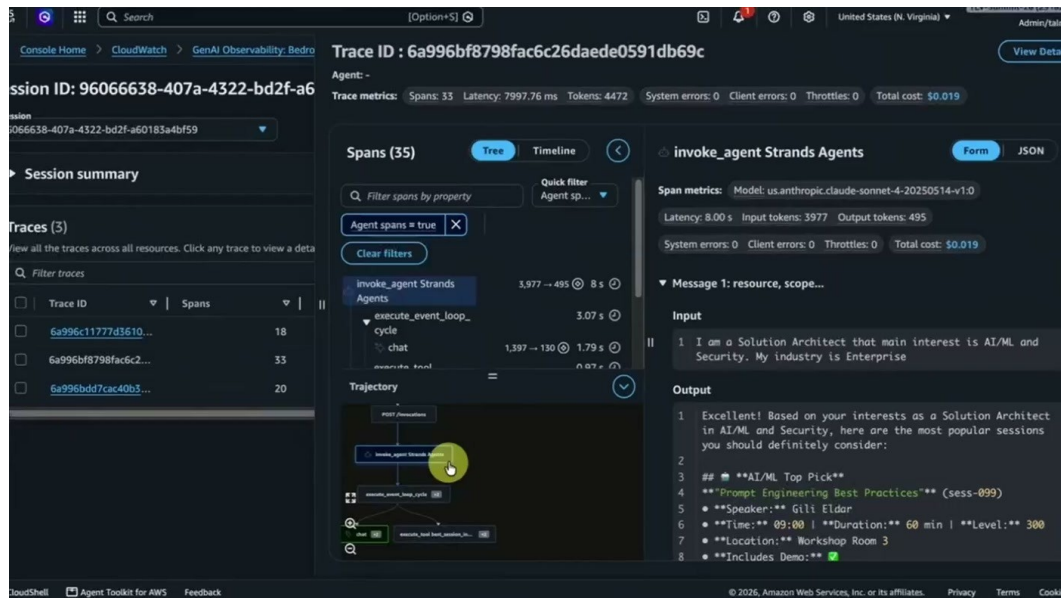
לאג'נט הוגדרו שלושה כלים: "יש לנו טול של Best Sessions in | Count Attendees, Suggest Sessions של Domain" [21:32]. עליהם הורצו שלוש בדיקות שונות, אחת מכל סוג — וזו הנקודה שבה החלק התיאורטי והדמו נפגשים.

בדיקה	סוג	מה היא קובעת
Tool Selection Accuracy	LLM-as-a-Judge	שהכלי שנבחר הוא הכלי המתאים; ציון בסקאלה 0-1 ועובר/לא עובר
Tool Parameter Accuracy	LLM-as-a-Judge	שהפרמטרים שעברו לכלי הם הנכונים, ושבדרך לא חלה טעות
Trajectory Exact Order Match	דטרמיניסטית	שהאג'נט השתמש בכלים בדיוק בסדר שהוגדר; "אם הוא יסתה טיפונת, זה אומר שהוא נכשל" [21:32]
Custom Evaluator	LLM-as-a-Judge מונחה	התנהגות כללית, כשלא יודעים את ה-trajectory המדויק

ל-Trajectory Exact Order Match נדרש להגדיר מראש את המסלול הצפוי דרך API, מול שירות ה-Ground Truth [21:32]. ה-Custom Evaluator, לעומת זאת, הוגדר כדי לבדוק "אם ההתנהגות הכללית נעשתה בדיוק כמו שביקשתי, למרות שאני לא יודע את הטראג'קטוריה המדויק" [23:03] — ההוראה שניתנה לו בדמו הייתה "based on the conversation context, whether the agent selected and used the correct tools" [27:47].

9.2 מה רואים ב-CloudWatch

בדמו עצמו נשאל האג'נט על ששנים מומלצים, ביקש רקע על המשתמש ("אני ארכיטקט שמעניין אותי בעצם דברים בסגנון ה-AI ML ואני נמצא באינדסטרי של אנטרפרייז"), החזיר המלצות, ואז קיבל סינון: "לא מתאים לי ששן שהוא לא לבל 300-400". משם עברו למסך.

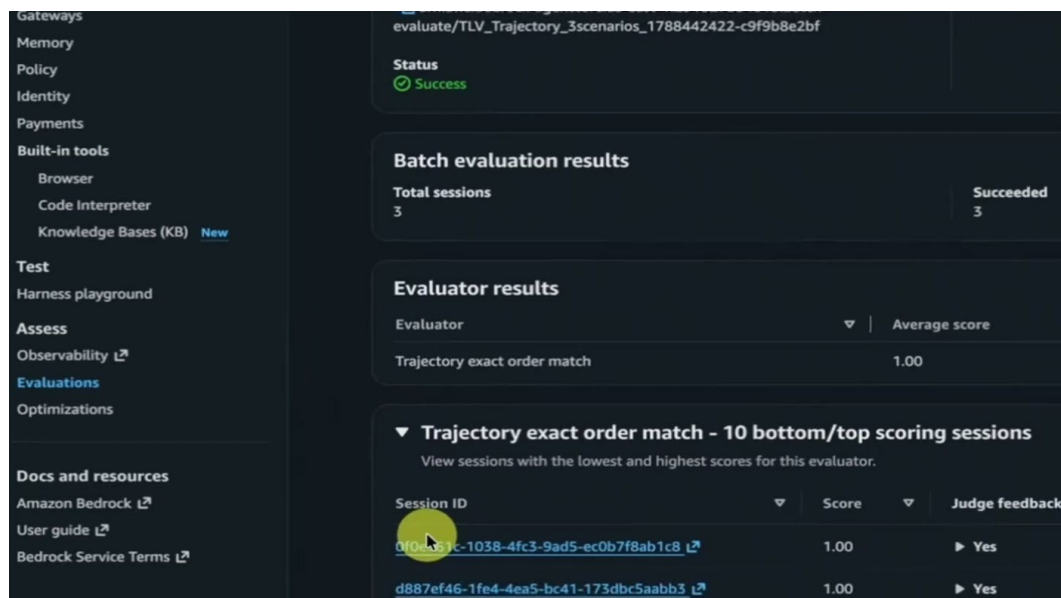


מסך ה-trace: רשימת ה-spans משמאל, ה-input וה-output של האג'נט מימין, וה-trajectory כציר יזואלי

לחיצה על trace חושפת "את כל האינפוט, את כל האוטפוט, את כל מה שביקשתי ממנו, את כל הריזנינג שהוא עשה" [23:03]. בצד השני של המסך מוצג ה-trajectory כציר — invoke agent, שימוש בכלי, לולאה חזרה — "את כל המידע הזה אני רואה בצורה ויזואלית" [24:42], עם זמני latency לכל span וגרף timeline.

— [23:03] עכשיו המידע הזה הוא אצלכם בתוך הקלאוד וואטש שלכם הוא לא יוצא החוצה

ה-evaluators הופעלו מתוך CloudWatch: נבחרו Tool Selection Accuracy ו-Parameter Accuracy, והוגדר sampling rate של 100% "כי אני רוצה את התוצאות כמה שיותר מהר" [23:03]. את התוצאות אפשר לראות ברמת אגרציה ואז לצלול פנימה — "התחלנו בגבוה ראינו את כלל המספרים וצללנו פנימה" [26:16].



תוצאות ה-batch evaluation: Total sessions מול Succeeded, וציון ממוצע 1.00 ל-Trajectory exact order match עם דירוג ה-sessions הגבוהים והנמוכים

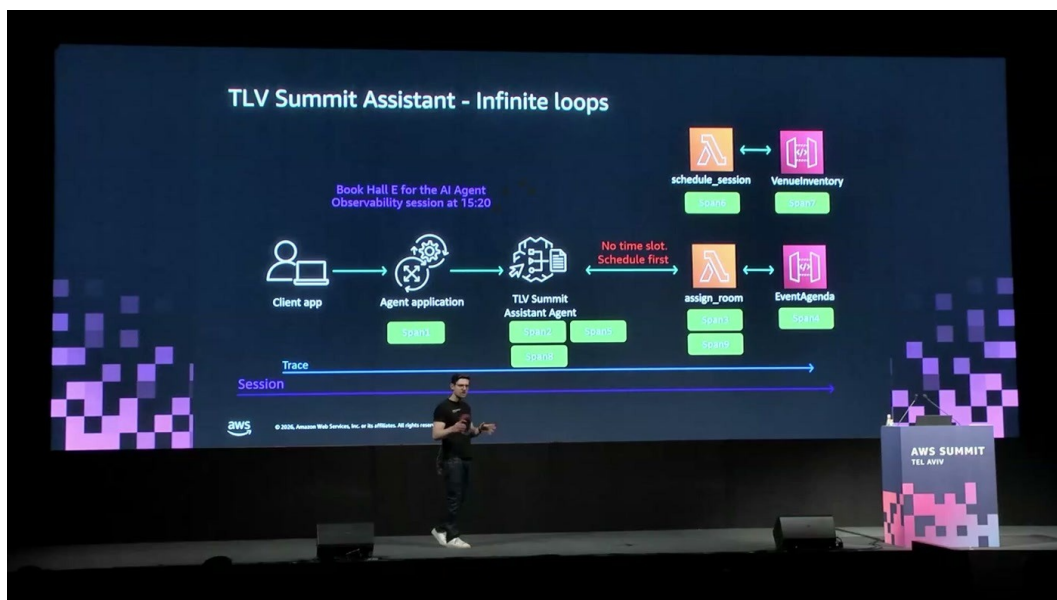
בבדיקה הדטרמיניסטית "ה-evaluation נתן סקור של average של 1", ופתיחת session בודד הראתה "yes, exact" [26:16] match actual trajectory — כלומר האג'נט עקב בדיוק אחרי המסלול שהוגדר לו מראש. ב-Custom Evaluator (שנקרא בדמו trajectory check online) התמונה הייתה מעניינת יותר: "היה לי פה 20 קאונטים של תשובות, 4 אירורים" [29:19], וגרף על ציר הזמן שמראה כמה תשובות היו נכונות וכמה לא.

הערת תמלול בהגדרת ה-Custom Evaluator נבחר מודל שיפוט מתוך רשימה, והוגדרה סקאלה עם שלוש נקודות — 1, 0 ו-0.5 — "כדי שאני אוכל לראות את היוריסטיקה ואת ההסתברויות" [27:47]. שם המודל שנבחר שובש בתמלול ולכן אינו מצוטט כאן.

הצעד האחרון בדמו הראשון היה ירידה ל-logs: דרך Log Analysis ב-CloudWatch נחקר ה-evaluation log, ובתוכו הופיע הנימוק של המודל — "התשובה של המודל מדוע הוא החליט לבחור את התשובה שהוא בחר" — לצד הציון עצמו, built-in tool selection accuracy עם ערך 1 [29:19]. ההיגיון שהודגש: ברמה הגבוהה רואים איפה הבעיה, ורק אז יורדים למיקרו. "הרי אם יש לי הרבה פעולות אני לא רוצה להיכנס אחד אחד ולהתחיל לחפש אותם" [29:19].

10. אתגר שני: Infinite Loops

האתגר השני הוגדר כך: "infinite loops זה מצב שהאג'נט שלנו לא מצליח לסיים משימה והוא נכנס ללולאה" [30:49]. הסיבות יכולות להיות רבות; הסשן צלל לאחת מהן דרך דוגמה שמדגימה אותה במדויק.



שני הכלים, assign_room ו-schedule_session, כשכל אחד מהם מתנה את פעולתו בכך שהשני כבר פעל

הפעם הפנייה לאג'נט הייתה מנקודת מבט של מארגני הסאמיט, לא של המשתתפים: "תסגור לי בבקשה את אולם E [...] לשעה שלוש ועשרים" [30:49]. מה שקרה:

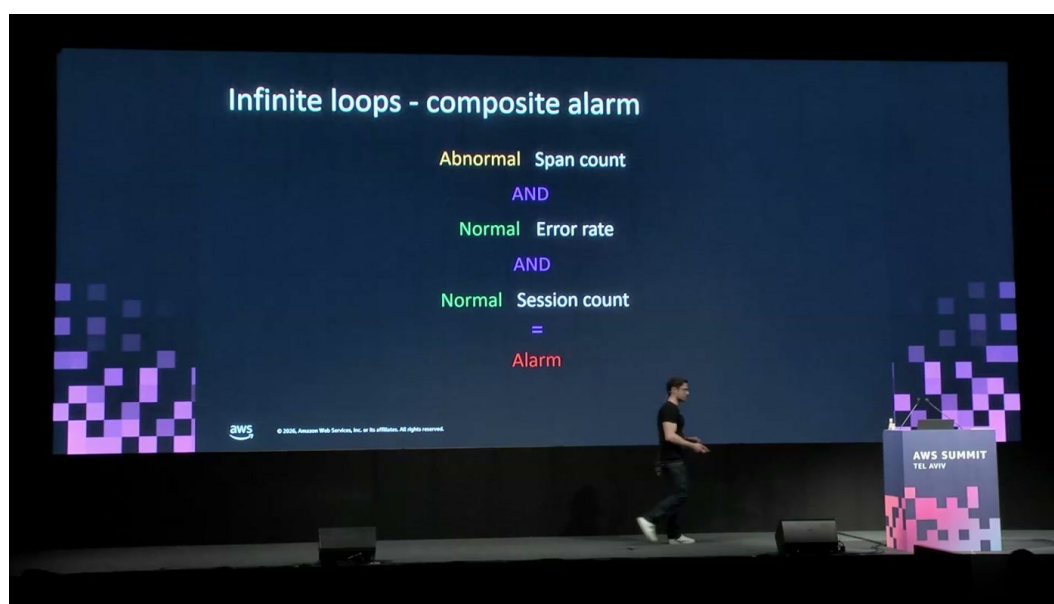
- **האג'נט פנה ל-assign_room.** assign_room בדק באג'נדה אם כבר זומנה הרצאה לשעה שלוש ועשרים, כי מבחינתו זה תנאי לשריון חדר. לא — ולכן החזיר: קודם תזמן את ההרצאה באג'נדה [32:21].
- **האג'נט פנה ל-schedule_session.** schedule_session בדק ב-inventory אם כבר שוריין חדר לאותה שעה, כי מבחינתו זה תנאי לזימון הרצאה. לא — ולכן החזיר: קודם תשריין חדר [32:21].
- **האג'נט חזר ל-assign_room עם אותה בקשה.** "וחוזר חלילה, אנחנו נמצאים בתוך לולאה" [32:21]. אף אחד משני הכלים לא נכשל; כל אחד מהם עשה בדיוק את מה שהוגדר לו.

סיוג המרצה הדגיש שזו דוגמה מאוד פשוטה, שנועדה להמחיש בעיה שרואים היום הרבה בתעשייה [32:21]. שימו לב לתכונה המשותפת לשני האתגרים בסשן: בשניהם המערכת עובדת בדיוק לפי ההגדרות, ובשניהם אין שגיאה שאפשר להתריע עליה.

11. זיהוי לולאה בסקייל: Composite Alarm

הדוגמה הראשונה הראתה לולאה ברמת session בודד, אבל "כשאנחנו מטפלים מערכת לא מעניין אותנו רזולוציה של session [...] מעניין אותנו רזולוציה של [33:55] scale. לשם כך הוצגו שלוש מטריקות שביחד מייצרות חתימה חד-משמעית.

מטריקה	מה היא סופרת	מה מצפים לראות בלולאה
Span count	כמה תחנות האג'נט עבר בהן בנקודת זמן	אנומליה — גבוה משמעותית מהרגיל
Error rate	כמה שגיאות היו באותה נקודת זמן	נורמלי לגמרי
Session count	כמה sessions עשו אינטראקציה עם האג'נט	נורמלי לגמרי



התנאי כפי שנוסח על הסלייד: Abnormal Span count AND Normal Error rate AND Normal Session count = Alarm

הקורלציה הזאת מיושמת ב-CloudWatch Composite Alarm — "פיצ'ר שיודע לקחת מספר התראות לבדוק אם מתקיים ביניהם איזשהו תנאי ואם הוא מתקיים להטריע" [33:55]. ההיגיון של שלילת החלופות הוסבר במפורש: קפיצה ב-error rate "יכלה להצביע על המון דברים [...] על בעיה באחת התחנות בדרך, על בעיה בכלי", וקפיצה ב-session count "יכלה להצביע אולי על שימוש חיובי במערכת שלנו" [35:26]. רק הצירוף מצביע על לולאה.

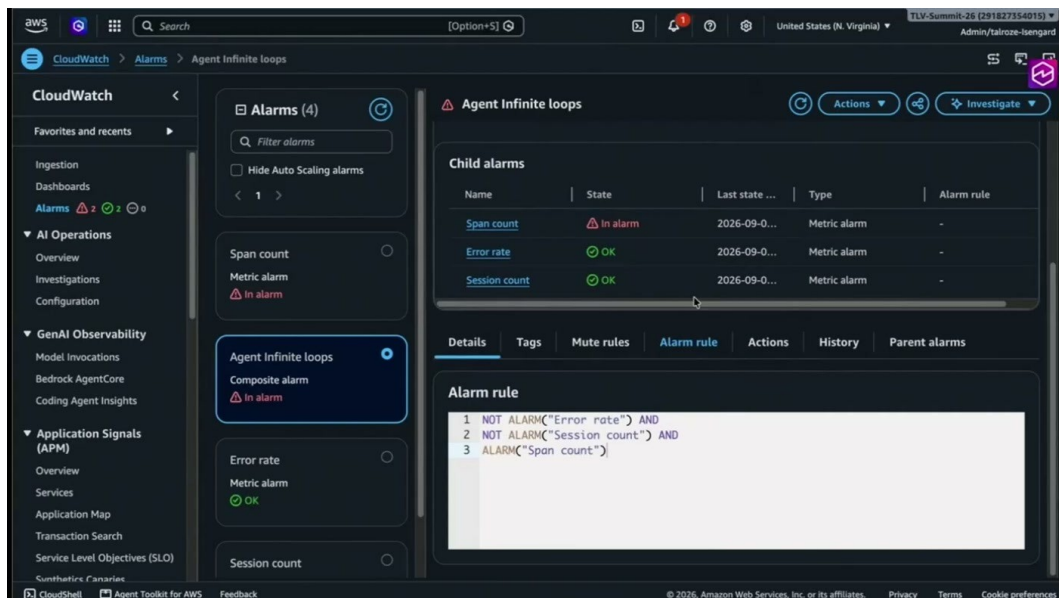
— [35:26] ברגע ששתנו את השילוב הזה שהכל רגיל מבחינת סשנים הכל רגיל מבחינת אירורים ואנחנו רואים קפיצה מאוד משמעותית במספר התחנות שהאג'ן שלנו עבר בהן אנחנו הולכים לדעת די בוודאות שאנחנו נמצאים באירוע של אינפנט לופ

11.1 איך בונים מטריקת span count מתוך לוג

הטריק המעשי של הדמו השני: מטריקת span count לא הייתה קיימת, אז היא נוצרה מתוך הלוגים. ב-CloudWatch Log Analytics נבנה query שמחפש את המילה span בלוגים של AgentCore, והמונה שלו הפך למטריקה. Rosenstein ציין שהוא לא כתב את ה-SQL בעצמו:

— [35:26] אז ביקשתי עם ה-LLM של CloudWatch לייצר לי query על Agent Core Span איפה שהוא מוצא את המילה Span שייצר לי בעצם query

הגרף שהתקבל הראה את כמות המופעים של המילה span ברגע נתון — "פה רואים 1440" [35:26] — וההערה הכללית שנלוותה: "הייתי יכול לבחור כל מילה אחרת ולייצר בעצם מטריקה מתוך מילה בלוג" [35:26]. על המטריקה הזאת נבנה alarm מסוג anomaly detection: הרצועה האפורה היא ה-band הצפוי, הקו הכחול הוא הקריאות בפועל, וההתראה מוגדרת על כל מה שמעל ה-band. [36:57]



ה-Composite Alarm במצב In alarm: שני ה-Error rate, Session count, ו-child alarms — OK — ורק Span count ב-alarm

שלושת ה-alarms אוחדו ל-Composite Alarm אחד לפי הכלל NOT ALARM("Error rate") AND NOT ALARM("Session count") AND ALARM("Span count"). כשהוזרמו נתונים, "שניים מתוך הבדיקות ציילד שלו שהוא עשה, ה-error rate וה-session count הם בירוק אבל הספן קאוונט עלה, אז הלארם הזה קפץ" [38:27]. הסלייד שקדם לדמו ציין ש-Composite Alarm בודד יכול לנטר עד 100 alarms בסיסיים — מנגנון להפחתת רעש התראות ברמת האגרגציה.

12. מה לוקחים מכאן

המרצים סגרו בהבהרה מה לא הייתה מטרת ההרצאה: "להציג לכם את האתגרים האלו ממש לא הייתה מטרת ההרצאה שלנו היום מטרת ההרצאה שלנו היום הייתה להציג לכם את צורת החשיבה האופרטיבית החדשה ואת הכלים" [38:27]. שני האתגרים היו רק הדגמה.

Key takeaways

- Agents bring new challenges - demanding a new operational mindset
- GenAI doesn't replace classic observability - it extends it
- Know where each one belongs

© 2026, Amazon Web Services, Inc. or its affiliates. All rights reserved.

שתי השורות שנשארו על המסך בסיום

- המבחן החדש הוא לא "האם יש שגיאה" אלא "האם המסלול נכון". שני האתגרים שהוצגו — tool misuse ולולאה אינסופית — קרו בלי שגיאה אחת. מערכת ניטור שמתרעה רק על כשלים לא תראה אף אחד מהם.
- **Evaluations** הוא כלי ניטור, לא רק כלי פיתוח. ההבחנה בין on-demand (פיתוח ו-CI/CD) ל-online (פרודקשן, קלטים חיים) היא ההבדל בין לבדוק שהאג'נט תקין פעם אחת ובין לדעת שהוא נשאר תקין [18:30].

- **בחרו את סוג הבדיקה לפי סוג השאלה.** LLM-as-a-Judge לשאלות איכות שאין להן תשובה מדויקת; code-based ו-trajectory match לשאלות שיש להן תשובה אחת. הדמו הריץ את שתיהן זו לצד זו על אותו אג'נט.
 - **הכלים הקלאסיים לא הלכו לשום מקום.** זיהוי הלולאה נעשה כולו ב-CloudWatch: מטריקה שנגזרה ממילא בלוג, anomaly detection, ו-composite alarm. שום דבר מזה אינו ייחודי ל-GenAI.
 - **הידע שצריך לרכוש הוא מתי להשתמש במה.** "תמיד תחשבו איפה אתם יכולים להשתמש בעולם הוא אבזרבביליטי הקלאסי [...]. ואיפה אתם יכולים להשתמש בעולם האבזרבביליטי החדש" [38:27].
 - **האימוץ לא מחייב לשנות framework.** AgentCore Runtime הוצג כ-framework agnostic, וה-observability מיוצאת ב-OpenTelemetry — כלומר גם למי שכבר עובד עם כלי ניטור אחר יש נתיב קליטה [20:02].
- בסלייד המשאבים שהוצג לקראת הסוף הופיעו שלושה QR codes: Generative AI observability, Debugging production agents with Amazon, Documentation, Evaluation Strategies for AI Agents, ו-Bedrock AgentCore Observability.

13. מילון מונחים

מונח	מה זה, כפי שהוגדר בסשן
Session / Trace / Span	שלוש רמות תיעוד: השיחה המלאה, אינטראקציה בודדת בתוכה, וצעד יחיד בתוך האינטראקציה
Trace context	"יומן המסלול" — המבנה שבו כל תחנה מתעדת את ה-ID שלה כך שאפשר לשחזר את הרצף
Trajectory	הנתיב שהאג'נט עשה בפועל בין הכלים; נבדק כ-exact order, in order או any order
Tool Misuse	האג'נט פנה לכלי לא נכון והחזיר תשובה לא נכונה — בלי שגיאה ובלי אינדיקציה
Infinite Loop	האג'נט לא מצליח לסיים משימה ונכנס ללולאה, למשל בין שני כלים שכל אחד מתנה את השני
AgentCore Runtime	הרצה serverless של agents, עם sessions מבודדים ומקבילות
AgentCore Memory	שמירת העדפות והקשר של המשתמש בין אינטראקציות
AgentCore Gateway	חשיפת כלים לאג'נט בצורה סטנדרטית, דרך פרוטוקול MCP
AgentCore Observability	איסוף traces, logs, spans ומטריקות מרכיבי AgentCore ושליחתם החוצה
AgentCore Evaluations	הרצת בדיקות מתוזמנות על פלט האג'נט בשלוש הרמות, עם ציון לכל בדיקה
LLM-as-a-Judge	מודל חיצוני שמעריך את פלט האג'נט, מנמק ונותן ציון בסקאלה 0–1
Custom Evaluator	בדיקה שכותבים בעצמכם ל-use case ספציפי, בכל אחת משלוש הרמות
Ground Truth	השירות שמולו מגדירים דרך API את המסלול הדטרמיניסטי הצפוי
OpenTelemetry (OTEL)	פרוטוקול טלמטריה מפרויקט CNCF, שדרכו הנתונים מיוצאים מהאג'נט החוצה
Composite Alarm	התראה ב-CloudWatch שמחברת מספר alarms בתנאי לוגי; יכולה לנטר עד 100 alarms בסיסיים