

# Cybersecurity and Law

2025 Nr 1(13)

DOI: 10.35467/cal/214578



## Comparison of Capabilities of Modern LLM Models in Analyzing and Developing National Security Strategies

### Porównanie współczesnych modeli LLM jako narzędzia wparcia analizy i tworzenia strategii cyberbezpieczeństwa państwa

**Tomasz ZAWADZKI**

Politechnika Warszawska, Wydział Mechatroniki

ORCID: 0000-0003-2758-6454

E-mail: tomasz.zawadzki2@pw.edu.pl

#### Abstract

The aim of the research described in this article was to assess the feasibility of using offline Large Language Models (LLMs) to analyze contemporary cybersecurity strategies of individual countries, identify their strengths and weaknesses, and subsequently prepare recommendations for designing national strategies. The research methodology adopted by the author first involved setting up the research environment, which included configuring hardware and software to enable efficient work with LLMs, followed by interacting with the models through appropriate queries and analyzing their responses. Depending on the answers provided, follow-up questions were asked for clarification. Importantly, the interaction was conducted in both Polish and English at the author's discretion, without a predefined pattern. Three of the most popular offline models were selected for interaction: LLaMA, QwQ, and DeepSeek-R1, each receiving the same set of questions. The outcome of these queries was a comparative analysis of the cybersecurity strategies of the USA, Germany, and Poland, as well as the strategy of the War Studies University, highlighting identified differences and similarities, and then preparing recommendations for a hypothetical Polish cybersecurity strategy for the years 2025–2035. The research also evaluated the models' familiarity with official national and local

documents and verified how these models perform processes of analysis, reasoning, and recommendation preparation, considering the governmental significance of the final document.

## **Keywords**

*LLaMA, QwQ, DeepSeek-R1, cybersecurity strategy*

## **Streszczenie**

Celem badań opisanych w niniejszym artykule była ocena możliwości wykorzystania dostępnych w wersji offline modeli LLM (ang. *Large Language Models*) do analizy współczesnych strategii cyberbezpieczeństwa poszczególnych państw, wyszukanie ich słabych i mocnych stron, a następnie przygotowanie rekomendacji w projektowaniu własnych strategii.

Przyjęta przez autora metodyka badawcza, w pierwszej kolejności polegała na przygotowaniu środowiska badawczego, czyli na odpowiedniej konfiguracji sprzętu i oprogramowania umożliwiającego wydajną pracę z modelami LLM, a następnie na interakcji z modelami za pomocą zadawania im odpowiednich zapytań oraz analizie odpowiedzi. W zależności od udzielonej odpowiedzi przesyłane były pytania uszczegółowiające. Co istotne, interakcja prowadzona była w języku polskim oraz angielskim wedle uznania autora, czyli bez określonego klucza. Do interakcji wybrane zostały 3 najpopularniejsze modele dostępne w trybie offline, tj. LLaMA, QwQ oraz DeepSeek-R1, a każdy z nich otrzymał taki sam zestaw pytań. Efektem zadanych pytań była analiza porównawcza strategii cyberbezpieczeństwa USA, Niemiec i Polski oraz dodatkowo strategii Akademii Sztuki Wojennej, uwypuklenie zidentyfikowanych różnic oraz podobieństw, a następnie przygotowanie rekomendacji dla hipotetycznie projektowanej strategii cyberbezpieczeństwa Polski na lata 2025-2035.

Efektem przeprowadzonych badań była ocena znajomości przez testowane modele oficjalnych dokumentów państwowych i lokalnych, a także weryfikacja w jaki sposób modele te przeprowadzają procesy analizy, wnioskowania oraz przygotowywania rekomendacji biorąc pod uwagę rangę państwową tworzonego finalnie dokumentu.

## **Słowa kluczowe**

*LLaMA, QwQ, DeepSeek-R1, strategia cyberbezpieczeństwa*

## **Wprowadzenie**

Prezentowane w niniejszym artykule wyniki stanowią kontynuację rozpoczętych w 2024 roku<sup>1</sup> przez autora badań, których głównym celem była identyfikacja kluczowych trendów w obszarze strategii cyberbezpieczeństwa państwa oraz ocena możliwości wykorzystania sztucznej inteligencji jako

---

<sup>1</sup> T. Zawadzki, *Sztuczna inteligencja jako narzędzie wsparcia w przygotowaniu strategii bezpieczeństwa państwa, Nowe horyzonty sztucznej inteligencji i uczenia maszynowego: trendy, wyzwania, perspektywy*, Warszawa 2024.

narzędzia wspierającego w procesie w jej tworzenia. Wówczas, autor skupił się głównie na ocenie potencjału modelu LLaMA 3.1 w wersji offline, jako że był on najbardziej upowszechniony i najlepiej wytrenowany, niemniej w związku gwałtownym rozwojem innych modeli, autor postanowił przeprowadzić drugą iterację badań, poszerzoną o porównanie dostarczanych przez nie wyników. Do badań wybrano najnowsze w wersji najpopularniejszych wówczas modeli tj. LLaMA 3.3, QwQ oraz DeepSeek-R1. Metodyka badań polegała, na konwersacji ze wskazanymi modelami na temat znajomości strategii cyberbezpieczeństwa wybranych krajów, a następnie przygotowaniu propozycji zapisów dla strategii Polski. W razie, uznania przez autora, iż wygenerowane odpowiedzi się są wystarczające wysyłane było żądanie o jej uszczegółowienie lub rozbudowanie. Jako podstawę do analizy wybrano obowiązujące strategię cyberbezpieczeństwa USA, Niemiec, Polski, czy Izraela, a także strategię rozwoju Akademii Sztuki Wojennej, jako dokument o znacznie mniejszej skali obowiązywania, jednak pokazujący, czy dany model posiada o nim wiedzę. Ponadto, aby odwzorować charakter „rozmowy” z modelami LLM, prowadzony naprzemiennie w języku polskim oraz angielskim, niektóre odpowiedzi są cytowane wprost, łącznie z wersją angielskojęzyczną. Efektem przeprowadzonego badań była ocena w jaki sposób wybrane do testów modele znają powszechnie obowiązują oraz publicznie dostępne dokumenty strategiczne, w jaki sposób potrafią je analizować oraz czy na podstawie posiadanej wiedzy są w stanie wygenerować nowe dokumenty o charakterze strategicznym, uwzględniając przy tym nowe zagrożenia, wynikające z chociażby szybkiego rozwoju sztucznej inteligencji. Należy także zwrócić uwagę, że do badania wybrane zostały wyłącznie modele, które umożliwiają pracę offline, gdyż autor zakłada, że przy wspierając się sztuczną inteligencją przy opracowywaniu dokumentów o charakterze strategicznym, ich twórcy nie podejmą ryzyka przekazywania kluczowych informacji o celach i sposobach funkcjonowania państwa poza wewnętrzne systemy teleinformatyczne.

### **LLaMA 3.3**

Model LLaMA 3.3 to najnowsza wersja dużego modelu językowego opracowanego przez firmę Meta, zaprezentowana 6 grudnia 2024 roku. Jest to model o 70 miliardach parametrów, zoptymalizowany pod kątem efektywności, wielojęzyczności i zastosowań praktycznych. W przeciwieństwie do swojego poprzednika, LLaMA 3.1, który występował również w wariantach 405B (405 miliardów parametrów), LLaMA 3.3 oferuje porównywalną wydajność przy znacznie niższych wymaganiach obliczeniowych. Choć jest ponad pięć razy mniejszy pod względem liczby parametrów, osiąga zbliżone wyniki w testach takich jak MMLU (Massive Multitask Language Understanding), HumanEval (kodowanie) czy IFEval (instruction following). Wykorzystuje mechanizm Grouped Query Attention (GQA) i obsługuje kontekst do 128 tys. tokenów.

LLaMA 3.3, w stosunku do badanej w pierwszej iteracji 3.1 wykazuje większy potencjał w zastosowaniach strategicznych, w tym w obszarze cyberbezpieczeństwa. Jego zdolność do przetwarzania języka naturalnego, rozumienia kontekstu oraz generowania logicznych wniosków pozwala na lepsze wykorzystanie w analizie dokumentów strategicznych, ocenie ryzyk, a także w projektowaniu nowych polityk bezpieczeństwa.

Model LLaMA 3.3 deklaruje znajomość wielu dokumentów strategicznych, m.in.:

- „Strategia Rozwoju Akademii Sztuki Wojennej na lata 2024–2035” – „Jest to strategiczny plan rozwoju Akademii Sztuki Wojennej w Warszawie, który określa cele i kierunki rozwoju uczelni na najbliższe lata.”
- „National Cybersecurity Strategy for USA (2023)” – „In March 2023, the Biden Administration released the National Cybersecurity Strategy, which outlines a new approach to cybersecurity”;
- „Cybersecurity Strategy for Germany 2021” – „In September 2021, the German Federal Government released the Cyber Security Strategy for Germany 2021”;
- „Strategia Cyberbezpieczeństwa RP na lata 2019–2024” – „The Cyber Security Strategy of the Republic of Poland for the years 2019–2024 is a comprehensive plan aimed at protecting Poland's digital infrastructure”.

Jednakże model przyznaje, że nie posiada bezpośredniego dostępu do treści tych dokumentów i bazuje na ogólnej wiedzy oraz wzorcach z danych treningowych.

Model potrafi przeprowadzić logiczną analizę wewnętrznej spójności strategii. W przypadku Strategii Rozwoju Akademii Sztuki Wojennej wskazuje:

- „Czy dokument określa jasno i precyzyjnie cel strategiczny Akademii Sztuki Wojennej? Czy cel ten jest realistyczny i możliwy do osiągnięcia?”;
- „Czy misja i wizja uczelni są spójne z celem strategicznym?”;
- „Czy dokument określa konkretne zadania i działania, które będą realizowane w celu osiągnięcia celu strategicznego?”;

Model LLaMA 3.3 przedstawił syntetyczne porównanie trzech strategii oraz wskazał kluczowe cele, na które są one nastawione:

1 USA:

- a. „Defending Critical Infrastructure”
- b. „Disrupting and Dismantling Threat Actors”
- c. „Building a More Resilient Digital Ecosystem”
- d. „Fostering a Global Cybersecurity Cooperation”

2. Niemcy:

- a. „Strengthening Cyber Security”
- b. „Improving Incident Response”
- c. „Enhancing International Cooperation”
- d. „Promoting Cyber Security Awareness”

3. Polska:

- a. „Improving Cyber Security”
- b. „Developing Cyber Security Capabilities”
- c. „Enhancing International Cooperation”
- d. „Promoting Cyber Security Awareness”.

Model wskazał również na słabości każdej strategii, np. „Lack of clear prioritization” w przypadku strategii polskiej, „Overemphasis on regulation” w strategii niemieckiej oraz „Overly broad scope” w strategii amerykańskiej.

Model wygenerował szczegółową propozycję strategii, zatytułowaną: „Cybersecurity Strategy for Poland 2025–2035: Protecting Digital Sovereignty and Fostering Innovation”. W strategii uwzględnił następujące elementy:

1. Cele strategiczne obejmujące:
  - a. Utworzenie Narodowej Agencji Cyberbezpieczeństwa
  - b. Ochrona infrastruktury krytycznej
  - c. Budowanie kultury cyberbezpieczeństwa
  - d. Rozwój sektora innowacji
  - e. Współpraca międzynarodowa
2. Elementy dodatkowe:
  - a. „Establish an AI Security Task Force”
  - b. „Develop Quantum-Resistant Cryptography Guidelines”
  - c. „Invest in Quantum Computing Research”
3. Ocena słabości i ryzyk strategii. Model wskazał na potencjalne słabości:
  - a. „Overemphasis on Technology”
  - b. „Lack of Clear Metrics”
  - c. „Insufficient Emphasis on Incident Response”
  - d. „Limited International Cooperation”
  - e. „No Clear Plan for Cybersecurity Workforce Development”
  - f. Dodatkowe środki zaradcze, takie jak: „Regular monitoring and evaluation”, „Diversified international partnerships”, „Dedicated budget allocation”
4. Budżet i zasoby ludzkie. Model oszacował budżet wdrożenia strategii na lata 2025–2035 na około 293 mln euro, z podziałem rocznym i sektorowym:
  - a. „National Cybersecurity Agency: 20%”
  - b. „Critical Infrastructure Protection: 25%”
  - c. „Cybersecurity Culture and Awareness: 15%”
  - d. „Cybersecurity Industry Development: 20%”
  - e. „International Cooperation: 10%”
  - f. Szacowana liczba osób zaangażowanych w realizację strategii to ok. 100–250 specjalistów, w tym zespoły ds. AI, kryptografii, SOC, edukacji i planowania.

Podsumowując, model LLAMA 3.3 wykazuje wysoką zdolność do:

- Analizy dokumentów strategicznych;
- Generowania propozycji polityk publicznych;

- Oceny ryzyk i słabości;
- Szacowania zasobów i kosztów.

Jednocześnie sam wskazał na swoje wewnętrzne ograniczenia związane z brakiem dostępu do pełnych treści dokumentów źródłowych. Powstaje więc pytanie, które odpowiedzi zostały wygenerowane na podstawie posiadanej wiedzy, a które tylko na podstawie np. tytułów dokumentów jak mogło to być w przypadku Strategii Akademii Sztuki Wojennej.

## Model QwQ

Model QwQ-32B, zwany powszechnie QwQ, to zaawansowany, średniej wielkości model językowy opracowany przez zespół Qwen w Alibaba Cloud, zaprezentowany w marcu 2025 roku jako część serii Qwen. Jego głównym celem jest rozwiązywanie złożonych problemów wymagających wieloetapowego rozumowania, takich jak zadania matematyczne, generowanie kodu, analiza logiczna czy ścisłe podążanie za instrukcjami. W odróżnieniu od klasycznych modeli czatujących, QwQ został zaprojektowany jako „model rozumujący”, który potrafi samodzielnie zadawać sobie pytania i weryfikować własne odpowiedzi w trakcie generacji. QwQ-32B posiada 32,5 miliarda parametrów i opiera się na głębokiej architekturze transformera z 64 warstwami. Wykorzystuje nowoczesne komponenty, takie jak SwiGLU, RMSNorm, Rotary Positional Embeddings (RoPE) oraz mechanizm Generalized Query Attention (GQA). Obsługuje kontekst o długości aż 131 072 tokenów, co pozwala mu analizować bardzo długie dokumenty lub prowadzić rozbudowane rozmowy bez utraty wcześniejszych informacji.

W toku prowadzonych badań modelu qwq, wykazywał się on znajomością ogólnych założeń strategii cyberbezpieczeństwa dla USA i Niemiec, ale nie posiada szczegółowych informacji o polskiej strategii na lata 2019–2024. Podsumowując, model udzielił następujących odpowiedzi:

- w przypadku strategii amerykańskiej model opisuje ją jako strategię skoncentrowaną na „defending critical infrastructure”, „disrupting cyber criminals”, „strengthening international partnerships” oraz „investing in R&D”. To pokazuje, że model zna strukturę i cele dokumentu, ale nie cytuje go bezpośrednio;
- w przypadku strategii niemieckiej model wskazuje na „enhancing cyber defense capabilities”, „incident response mechanisms”, „public-private partnerships” i „education and research”. To są typowe komponenty strategii UE;
- w przypadku strategii polskiej model przyznaje, że „I don't have specific information about the Polish cybersecurity strategy for the years 2019–2024.” To oznacza, że nie został wytrenowany na pełnym tekście tego dokumentu lub nie ma do niego dostępu.

Na podstawie uzyskanych odpowiedzi, można wnioskować, że model nie generuje fałszywych i niepotwierdzonych treści, co jest zgodne z zasadami bezpieczeństwa i transparentności, a jego odpowiedzi są oparte na znanych wzorcach.

Model qwq zapytany o strukturę strategii cyberbezpieczeństwa przedstawił 10 kluczowych komponentów, które powinny znaleźć się w takim dokumencie i uzasadnił konieczność ich zawarcia. Poniżej przedstawione zostały wskazane komponenty w raz z uzasadnieniem:

1. Executive Summary – streszczenie celu strategii, zakresu i kluczowych interesariuszy.
  - a) Uzasadnienie: pozwala decydom szybko zrozumieć istotę dokumentu.
2. Risk Assessment – identyfikacja zagrożeń, analiza podatności, ocena ryzyka.
  - a) Uzasadnienie: umożliwia priorytetyzację działań i alokację zasobów.
3. Policy and Compliance – polityki bezpieczeństwa, zgodność z regulacjami (np. GDPR, NIS).
  - a) Uzasadnienie: zapewnia legalność działań i ochronę danych.
4. Governance and Roles\*\* – struktura organizacyjna, przypisanie ról i odpowiedzialności.
  - a) Uzasadnienie: zapobiega chaosowi i umożliwia skuteczne zarządzanie.
5. Protection Measures – środki techniczne (firewalle, IDS, szyfrowanie)
  - a) Uzasadnienie: to fundament obrony przed atakami.
6. Detection and Response – monitoring, SIEM, plan reagowania.
  - a) Uzasadnienie: szybka reakcja minimalizuje skutki incydentów.
7. Recovery Plan – kopie zapasowe, plan ciągłości działania.
  - a) Uzasadnienie: zapewnia odporność i zdolność do powrotu do działania.
8. Training and Awareness – szkolenia, kampanie edukacyjne.
  - a) Uzasadnienie: człowiek jest najsłabszym ogniwem – edukacja zmniejsza ryzyko.
9. Vendor Management – ocena dostawców, umowy, monitoring.
  - a) Uzasadnienie: łańcuch dostaw to częsty wektor ataku.
10. Continuous Improvement – cykliczne przeglądy, aktualizacje, R&D.
  - a) Uzasadnienie: cyberzagrożenia ewoluują – strategia musi być dynamiczna.

Dodatkowo, w czasie dyskusji model qwq wskazał na nowe zagrożenia, związane ze sztuczną inteligencją (AI) i podał ich przykłady:

- Automatyzacja ataków – AI może generować phishing na masową skalę,
- Adversarial ML – techniki omijania systemów wykrywania,
- Deepfake – manipulacja obrazem i dźwiękiem w celu oszustwa,
- Predykcja podatności – AI może analizować systemy i wskazywać słabe punkty,
- Botnety oparte na AI – inteligentne, rozproszone ataki,

a następnie zaproponował konkretne środki zaradcze, takie jak:

- Inwestycje w „defensywne AI” – systemy uczące się na bieżąco,
- Red teaming z AI – symulacja ataków w celu testowania obrony,
- Ramy regulacyjne – kontrola nad wykorzystaniem AI w cyberbezpieczeństwie,
- Współpraca międzynarodowa – wymiana wiedzy i wspólne projekty.

Zapytany z kolei o budżet wdrożenia strategii w zaproponowanym przez siebie kształcie przedstawił szczegółowy plan finansowy, który uwzględnia zarówno podstawowe komponenty strategii, jak i dodatkowe środki na przeciwdziałanie zagrożeniom AI. Budżet ten podzielony był na trzy części:

1. Podstawowy budżet (bez AI):
  - a) Wartość całkowita: 1 415, 64 mln EUR w ciągu 10 lat,
  - b) Roczna alokacja: 141 564 mln EUR.
2. Dodatkowy budżet na AI:
  - a) R&D: 100 mln EUR,
  - b) Szkolenia: 50 mln EUR,
  - c) Technologie AI: 150 mln EUR,
  - d) Regulacje: 20 mln EUR,
  - e) Współpraca międzynarodowa: 30 mln EUR,
  - f) Łącznie: 350 mln EUR.
3. Łączny budżet:
  - a) 1 765,64 mln EUR w ciągu 10 lat,
  - b) Rocznie: 176 564 mln EUR.

Model wskazał również szczegółowy podział roczny na kategorie, takie jak:

- a) Personel: 60 mln EUR,
- b) Technologie: 50 mln EUR,
- c) R&D: 20 mln EUR,
- d) Szkolenia: 15 mln EUR,
- e) Regulacje: 5 mln EUR,
- f) Współpraca: 6,564 mln EUR,

Na pytanie o porównanie strategii cyberbezpieczeństwa Polski ze jej odpowiednikami amerykańskim oraz niemieckim, model wskazał odpowiednie:

1. Cechy wspólne
  - a) Ochrona infrastruktury krytycznej
  - b) Inwestycje w badania i rozwój
  - c) Współpraca międzynarodowa
  - d) Ramy prawne i zgodność z regulacjami
2. Różnice
  - a) Polska kładzie większy nacisk na zagrożenia regionalne (np. ze strony Rosji),
  - b) Silna integracja z polityką UE,
  - c) Wzmacnianie zdolności obronnych w ramach NATO,
  - d) Udział w programach takich jak „Eastern Partnership”.

Model podkreśla, że strategia Polski jest dostosowana do jej geopolitycznego położenia, potrzeb narodowych oraz struktur europejskich\*\*.

Biorąc powyższe pod uwagę model qwq wykazuje wysoką zdolność do:

1. Strukturalnej analizy strategii cyberbezpieczeństwa;
2. Generowania realistycznych komponentów strategii;
3. Uwzględniania kontekstu regionalnego (Polska, UE);
4. Rozpoznawania zagrożeń AI i integracji ich w strategię;
5. Szacowania budżetu i alokacji zasobów.

Jednocześnie model zachowuje ostrożność w ocenie politycznej i nie podejmuje się krytyki strategii narodowych, co jest zgodne z zasadami bezpieczeństwa.

## DeepSeek-R1

Model DeepSeek-R1 to zaawansowany, otwartoźródłowy model językowy opracowany przez chińską firmę DeepSeek AI i wydany w styczniu 2025 roku. Jego głównym celem jest rozumowanie, rozwiązywanie problemów matematycznych, generowanie kodu oraz analiza logiczna. W odróżnieniu od klasycznych modeli opartych na gęstych transformatorach, DeepSeek-R1 wykorzystuje architekturę Mixture-of-Experts (MoE), która pozwala na aktywację tylko najbardziej odpowiednich „ekspertów” dla danego zadania. Dzięki temu, model osiąga wysoką wydajność przy znacznie niższym koszcie obliczeniowym – mimo że posiada 671 miliardów parametrów, podczas pracy aktywowanych jest tylko około 37 miliardów. Obsługuje kontekst o długości 128 tys. tokenów, co pozwala na analizę bardzo długich dokumentów i prowadzenie rozbudowanych rozmów, a w procesie treningu zastosowano unikalne podejście: reinforcement learning (RL) z etapami samodzielnego odkrywania wzorców rozumowania i ich dostosowywania do preferencji użytkowników.

W przypadku interakcji z modelem Deepseek-R1 wielokrotnie zaznaczał on, że nie zna języka polskiego i nie ma bezpośredniego dostępu do dokumentów źródłowych (np. strategii USA, Polski czy Niemiec), ale mimo to generuje logiczne, rozbudowane analizy. Przykładowo, przy pytaniu o strategię Akademii Sztuki Wojennej, model samodzielnie interpretuje nazwę i tworzy strukturę strategiczną, opartą na ogólnych założeniach edukacyjnych i organizacyjnych. Model „udając”, że zna ten dokument błędnie tłumaczy nazwę instytucji jako „Military Academy of Music”, ale mimo tego tworzy trafną analizę strategiczną. Wymienia jej następujące elementy:

1. Rozszerzenie programu nauczania,
2. Integrację technologii (AI, VR),
3. Współpracę międzynarodową,
4. Rozwój kadry i infrastruktury.

To pokazuje, że mimo błędnej interpretacji nazwy, model potrafi wygenerować realistyczne punkty strategii instytucjonalnej, niemniej łatwo

zauważyć, że są to bardzo ogólne założenia, które można dopasować do każdej jednostki edukacyjnej.

Model DeepSeek-R1 zapytanie o wiedzę o strategiach cyberbezpieczeństwa USA, Niemiec i Polski, wprost przyznaje, że nie zna treści dokumentów, ale tworzy ich logiczne struktury, tj.:

1. USA „The strategy likely includes goals such as defending critical infrastructure, promoting cybersecurity education”;
2. Niemcy: „Germany has been proactive in international cybersecurity initiatives”;
3. Polska: „While i don't have access to the exact content... i can provide general context”.

W każdym z tych przypadków model tworzył strukturę opartą o znane powszechnie komponentach strategii narodowych mimo, iż sam przyznał, że bezpośrednio nie zna tych dokumentów.

Zapytany o zaproponowanie nowej. Struktura strategii cyberbezpieczeństwa dla Polski generuje strukturę strategii, zawierającą następujące elementy:

1. Executive Summary,
2. Risk Assessment,
3. Action Plan,
4. Compliance,
5. Monitoring,
6. Incident Response,

przy czym dodatkowo sugeruje, aby uwzględnić zapisy dostosowaną do kontekstu Polski takie jak:

1. położenie geopolityczne,
2. zagrożenia ze strony państw trzecich,
3. wysoką zależność od usług cyfrowych,
4. potrzebę szyfrowania post-kwantowego,
5. współpracę z UE,

co pokazuje, że pomimo braku dostępu do istniejących dokumentów źródłowych potrafi wywnioskować jak dostosować strategię do specyfiki kraju w zależności od regionu świata, w którym jest on położony.

Model ten zaproponował także kluczowe wskaźniki efektywności (KPI), które miałyby pomóc w monitorowaniu stopnia wdrażania strategii. Zaproponowane wskaźniki to:

1. liczba naruszeń danych,
2. czas reakcji na incydenty,
3. poziom świadomości pracowników,
4. wskaźnik wykrywania prób ataków.

Deepseek-R1 oszacował także budżet i alokację środków na lata 2025–2035) na poziomie od 1,1 mld do 1,9 mld EUR w skali roku, z rosnącym udziałem AI (od 30% do 45%). Wskazał także, że budżet powinien uwzględniać następujące składowe:

1. Technologie,
2. Szkolenia,
3. Współpracę międzynarodową,
4. R&D,
5. Kampanie społeczne.

## Podsumowanie

Bazując na analizie odpowiedzi wygenerowanych przez poszczególne modele, można stwierdzić, że każdy z badanych modeli jest w stanie wnieść unikalną wartość do procesu tworzenia strategii, a ich połączenie może istotnie usprawnić opracowywanie kompleksowego dokumentu strategicznego dla Polski, uwzględniając, że;

- LLAMA 3.3 sprawdza się jako generator struktury dokumentu – idealny do tworzenia szablonów strategii,
- QwQ jest najbardziej zaawansowany w zakresie planowania strategicznego, budżetowania i analizy zagrożeń związanych z rozwojem AI,
- DeepSeek-R1 wykazuje dużą elastyczność i zdolność do interpretacji kontekstu, nawet przy ograniczonej znajomości języka i dokumentów źródłowych.

W Tabeli 1 przedstawiono syntetyczne porównanie działania zbadanych modeli w kontekście analizy i wsparcia przy tworzeniu strategii cyberbezpieczeństwa.

**Tabela 1.** Charakterystyka otrzymanych odpowiedzi przez poszczególne modele

| Kryterium                              | LLAMA                                                                 | QwQ                                                                         | DeepSeek-R1                                                                    |
|----------------------------------------|-----------------------------------------------------------------------|-----------------------------------------------------------------------------|--------------------------------------------------------------------------------|
| <b>Znajomość dokumentów źródłowych</b> | Nie zna konkretnych dokumentów, tworzy strukturę na podstawie wzorców | Nie zna treści strategii RP 2019–2024, ale tworzy realistyczne ramy         | Nie zna języka polskiego, ale dedukuje strukturę strategii z kontekstu         |
| <b>Zdolność do syntezy strategii</b>   | Bardzo dobra – tworzy 12-punktową strukturę strategii                 | Bardzo dobra – dodaje komponenty AI, budżet, KPI, porównania międzynarodowe | Bardzo dobra – rozbudowana analiza strategii edukacyjnej i cyberbezpieczeństwa |
| <b>Uwzględnienie AI w strategii</b>    | Wspomina o AI jako narzędziu wspierającym                             | Szczegółowo analizuje zagrożenia AI i proponuje środki zaradcze             | Rozpoznaje AI jako kluczowy element technologiczny i zagrożeniowy              |
| <b>Planowanie budżetowe</b>            | Nie podaje konkretnych kwot                                           | Szacuje budżet roczny (€1.1–1.9                                             | Potwierdza potrzebę inwestycji w AI, edukację, infrastrukturę i regulacje      |

| Kryterium                     | LLAMA                                              | QwQ                                                                  | DeepSeek-R1                                                              |
|-------------------------------|----------------------------------------------------|----------------------------------------------------------------------|--------------------------------------------------------------------------|
|                               |                                                    | mld), z rosnącym udziałem AI                                         |                                                                          |
| <b>KPI i mierniki sukcesu</b> | Wspomina ogólnie o ewaluacji                       | Proponuje 10 konkretnych wskaźników (np. czas reakcji, wykrywalność) | Wskazuje na potrzebę mierzenia efektów strategii (np. liczba partnerstw) |
| <b>Rekomendacje dla rządu</b> | Ogólne – wdrożenie struktury, edukacja, compliance | Szczegółowe – AI Task Force, regulacje, współpraca międzynarodowa    | Praktyczne – kampanie społeczne, rozwój kadry, integracja z UE/NATO      |
| <b>Harmonogram wdrożenia</b>  | Nie podaje                                         | Proponuje podział na 5 etapów (2025–2035)                            | Tworzy logiczny harmonogram wdrożenia strategii edukacyjnej i cyfrowej   |
| <b>Styl wypowiedzi</b>        | Formalny, uporządkowany, techniczny                | Analityczny, budżetowy, strategiczny                                 | Dedukcyjny, narracyjny, z próbą interpretacji kontekstu                  |

**Źródło:** opracowanie własne.

Z przedstawionego porównania wynika, iż każdy z modeli udziela inaczej zorientowanych odpowiedzi, co prawdopodobnie wynika z ich innych konstrukcji oraz założeń projektowych. Warto zwrócić uwagę na konstrukcję modelu DeepSeek-R1, który wprost przyznaje, że nie zna wskazanych przez autora dokumentów, niemniej zgodnie ze swoją konstrukcją próbuje wnioskować o ich naturze, na podstawie przekazanych informacji. Jest to istotna różnica w stosunku do pozostałych dwóch modeli, które swoje odpowiedzi starają się generować wyłącznie w oparciu o posiadaną wiedzę, natomiast nie dedukują tak głęboko, bazując tylko na fragmentach posiadanych informacji. Może to być zarówno zaleta jak i wada, gdyż z jednej strony model ten może wygenerować nieoczywiste wnioski i rekomendacje, a z drugiej mogą być one bardzo mało prawdopodobne. Niemniej w związku z ze stałym rozwojem modelu LLM ocena ich przydatności pod kątem wsparcia w przygotowywaniu dokumentów o charakterze strategicznym powinna być przedmiotem dalszych badań.

## Bibliografia

- A. Rahman, S.H. Mahir, Tanjum An Tashrif, etal., Comparative analysis based on DeepSeek, ChatGPT, and Google Gemini: Features, techniques, performance, future prospects, „Systems and Soft Computing” 2025, vol. 7.
- Artificial Intelligence and National Security, Center for Security and Emerging Technology, 2020.

- Pavic A., Berisa H., Artificial Intelligence and national security strategy development – challenges and perspectives, „National Interest” 2024, Rok XX, vol. 48, No. 2.
- Srivastava K., Artificial Intelligence and national security: perspective of the global south, „International Journal of Law in Changing World” 2023, vol. 2, Issue 2.
- Włodyka E. M., Artificial Intelligence as a Supporting Tool for Local Government Decision, „Defence Science Review” 2023, No. 17.
- Mikhailov D. I., Optimizing National Security Strategies through LLM-Driven Artificial Intelligence Integration.
- Zawadzki T., Sztuczna inteligencja jako narzędzie wsparcia w przygotowaniu strategii bezpieczeństwa państwa, Nowe horyzonty sztucznej inteligencji i uczenia maszynowego: trendy, wyzwania, perspektywy, Warszawa 2024.
- [https://www.researchgate.net/publication/370697446\\_Optimizing\\_National\\_Security\\_Strategies\\_through\\_LLM-Driven\\_Artificial\\_Intelligence\\_Integration](https://www.researchgate.net/publication/370697446_Optimizing_National_Security_Strategies_through_LLM-Driven_Artificial_Intelligence_Integration).
- [https://www.bbn.gov.pl/ftp/dokumenty/Strategia\\_Bezpieczenstwa\\_Narodowego\\_RP\\_2020.pdf](https://www.bbn.gov.pl/ftp/dokumenty/Strategia_Bezpieczenstwa_Narodowego_RP_2020.pdf).
- <https://www.gov.pl/web/cyber-nccpl/strategia-Cyberbezpieczenstwa-RP>.
- <https://aszwoj.bip.gov.pl/strategia-akademii-sztuki-wojennej/strategia-aszwoj.html>.
- <https://www.gov.pl/web/cyber-nccpl/strategia-Cyberbezpieczenstwa-RP>.
- <https://www.bmi.bund.de/EN/topics/it-internet-policy/cyber-security-strategy/cyber-security-strategy-node.html>.
- <https://www.whitehouse.gov/wp-content/uploads/2023/03/National-Cybersecurity-Strategy-2023.pdf>.
- <https://aszwoj.bip.gov.pl/fobjects/details/1747530/strategia-rozwoju-akademii-sztuki-wojennej-na-lata-2024-2035-pdf.html>.
- [https://owasp.org/www-project-top-10-for-large-language-model-applications/assets/PDF/OWASP-Top-10-for-LLMs-2023-v1\\_1.pdf](https://owasp.org/www-project-top-10-for-large-language-model-applications/assets/PDF/OWASP-Top-10-for-LLMs-2023-v1_1.pdf).
- <https://huggingface.co/Qwen/QwQ-32B>.
- <https://ollama.com/library/qwq>.
- <https://ollama.com/library/llama3.3>.
- <https://huggingface.co/meta-llama/Llama-3.3-70B-Instruct>.
- <https://github.com/deepseek-ai/DeepSeek-R1>.
- <https://ollama.com/library/deepseek-r1>.