

OPIS PRZEDMIOTU ZAMÓWIENIA

Przedmiotem zamówienia jest **dostawa infrastruktury obliczeniowej (zwanej dalej „Systemem”), oprogramowania wraz z wdrożeniem oraz świadczeniem usługi gwarancyjnej.**

1. Termin realizacji zamówienia:

Zgodnie ze złożoną ofertą nie później niż w terminie:

- do 60 dni od podpisania Umowy - dostarczenie Sprzętu;
- do 80 dni od podpisania Umowy - montaż, konfigurację i instalację dostarczonego Sprzętu wraz z oprogramowaniem;
- usługa gwarancyjna – do końca trwania okresu zobowiązania, tj. 60 miesięcy.

2. Zamówienie obejmuje:

- 2.1. Dostawę Sprzętu opisanego szczegółowo w pkt 2, na własny koszt, do jednego z ośrodków przetwarzania danych Zamawiającego na terenie m.st. Warszawy.
- 2.2. Montaż, konfigurację i instalację dostarczonego Systemu wraz z oprogramowaniem.
- 2.3. Przeprowadzenie instruktażu stanowiskowego dla personelu Zamawiającego w zakresie obsługi Systemu oraz oprogramowania.
- 2.4. Usługa gwarancyjna, w ramach której wchodzi zapewnienie wsparcia technicznego, usuwania awarii i błędów, oraz dostęp do najnowszych wersji oprogramowania w okresie min. 60 miesięcy od daty odbioru instalacji i konfiguracji – uruchomienia produkcyjnego Systemu.

3. Sprzęt - wymagania:

Przedmiotem zamówienia jest dostawa i wdrożenie wysokowydajnego Systemu, przeznaczonego do realizacji zadań z zakresu sztucznej inteligencji (AI, SI) i uczenia maszynowego (ML). System ma na celu zapewnienie skalowalnej i efektywnej platformy do trenowania i wdrażania zaawansowanych modeli AI, przetwarzania dużych zbiorów danych oraz prowadzenia badań naukowych.

Zamawiający wymaga dostarczenia kompletnego Systemu obejmującego serwery z zainstalowanymi akceleratorami GPU dedykowanymi do wsparcia tworzenia i trenowania modeli AI oraz uczenia maszynowego oraz przetwarzania dużych zbiorów danych, macierze dyskowe i przełączniki sieciowe w postaci standardowej, certyfikowanej przez producenta akceleratorów GPU konfiguracji, przy czym konfiguracja ta ma być zgodna z referencyjną konfiguracją producenta akceleratorów GPU o wymaganej skali i mocy obliczeniowej. Wykonawca, który zaoferuje System, w ramach którego producent akceleratorów GPU oraz przełączników sieci LAN i SAN będzie również producentem całego Systemu otrzyma dodatkowe punkty w kryterium Jakość. Zamawiający wymaga dostarczenia kompletnego, zintegrowanego przez jednego producenta Systemu z następujących powodów:

Integracja: Zapewnienie pełnej, niezawodnej oraz wydajnej integracji wszystkich komponentów sprzętowych oraz programowych eliminując problemy związane z ewentualnym brakiem kompatybilności.

Usługa gwarancyjna obejmująca wsparcie techniczne: W przypadku problemów technicznych, jeden producent systemu wraz z dostawcą zapewnią kompleksowe wsparcie, minimalizując czas potrzebny na ich rozwiązanie oraz eliminując problemy z ustaleniem odpowiedzialności.

Wydajność: System musi charakteryzować się wysoką wydajnością i skalowalnością, aby umożliwić Zamawiającemu wykonywanie rozproszonych obliczeń na wielu węzłach jednocześnie. Producent systemu musi gwarantować skalowalność i wydajność Sprzętu potwierdzoną niezależnymi testami.

Zamawiający wymaga, aby oferowany System był uwzględniony w rankingu MLCommons. Ranking MLCommons to uznawany w branży, transparenty i niezależny standard oceny wydajności sprzętu i oprogramowania w zadaniach związanych ze sztuczną inteligencją (AI) i uczeniem maszynowym (ML). Jego obecność w warunkach wyboru dostawcy ma na celu zagwarantowanie, że Zamawiający otrzyma System spełniający rygorystyczne standardy, optymalnie dopasowany do potrzeb Zamawiającego.

3.1. Wymagania ogólne

Wydajność: System musi charakteryzować się wysoką wydajnością w zakresie obliczeń numerycznych, operacji wejścia/wyjścia oraz komunikacji między węzłami, zgodnie ze specyfikacją i wymaganiami poniżej.

Skalowalność: System musi umożliwiać łatwą rozbudowę w przyszłości, zarówno pod względem liczby serwerów obliczeniowych, jak i pojemności macierzy dyskowych oraz przepustowości sieci.

Niezawodność: System musi charakteryzować się wysoką niezawodnością i dostępnością, minimalizując ryzyko przestojów i utraty danych.

Zarządzanie: System musi być wyposażony w intuicyjne i funkcjonalne narzędzia do zarządzania i monitorowania jego pracy.

Kompatybilność: System musi być kompatybilny z popularnymi frameworkami i bibliotekami wykorzystywanymi w obszarze AI/ML, w szczególności musi być w pełni kompatybilny z TensorFlow, PyTorch.

Kompletność: System musi być kompletny. Zamawiający oczekuje dostarczenia wszystkich elementów niezbędnych do instalacji, konfiguracji i uruchomienia Systemu. W szczególności Zamawiający oczekuje dostarczenia wraz z Systemem niezbędnego okablowania, obudów wraz z szynami montażowymi.

3.2. Wymagania szczegółowe

Lp.	Obszar	Wymaganie
1.	System	1.1. Oferowany System musi znajdować się na liście niezależnie zweryfikowanych Systemów przez MLCommons. 1.1.1. Oferowany System musi znajdować się na najnowszej na dzień publikacji Zapytania (Wersja 4.1 z dnia 17.01.2025) liście niezależnie zweryfikowanych Systemów przez MLCommons w zakresie trenowania modeli AI. Lista dostępna jest pod adresem https://mlcommons.org/benchmarks/training/, tabela MLPerf Results, benchmark: Training. 1.1.2. Zweryfikowana przez MLCommons konfiguracja oferowanego Systemu w zakresie trenowania modeli AI musi dotyczyć wielu węzłów (tj. dotyczyć Systemu składającego się z więcej niż jednego serwera wyposażonego w akceleratory GPU). 1.1.3. Zweryfikowana przez MLCommons konfiguracja oferowanego Systemu w zakresie trenowania modeli AI musi potwierdzać możliwość trenowania największego modelu językowego objętego weryfikacją – gpt3 (Kolumna gpt3 w tabeli MLPerfResults) 1.1.4. Oferowana konfiguracja Systemu musi składać się z co najmniej jednej zweryfikowanej przez MLCommons konfiguracji Systemu spełniającej powyższe wymagania, przy czym dopuszczalne jest zwielokrotnienie tej konfiguracji poprzez dostarczenie większej liczby węzłów. Zamawiający oczekuje, że system zostanie dostarczony w konfiguracji nie więcej niż 31 węzłów compute, zawierających każdy nie więcej niż 8 akceleratorów GPU wspierających obliczenia, dostarczonych w szafach rack wraz z niezbędnymi komponentami zarządzania systemem, dystrybucji zasilania, sieci LAN oraz sieci storage. 1.1.5. Oferowany System musi znajdować się na najnowszej liście (na dzień publikacji Zapytania – wersja 5.0 z dnia 25.03.2025) niezależnie zweryfikowanych przez MLCommons Systemów w zakresie uruchamiania modeli AI w centrach przetwarzania danych. Lista dostępna jest pod adresem https://mlcommons.org/benchmarks/inference-datacenter/, tabela MLPerf Results, benchmark: inference 1.1.5.1. Zweryfikowana przez MLCommons konfiguracja oferowanego Systemu w zakresie uruchamiania modeli AI musi prezentować wydajność uruchamiania dużych modeli językowych w architekturze MoE (Mixture of Experts) w postaci wyniku pomiaru wydajności uruchamiania na tej

		konfiguracji modelu gęstego llama3.1-405b (kolumna llama3.1-405b, Server Queries/s w tabeli Inference) oraz modelu eksperckiego mixtral-8x7b (kolumna mixtral-8x7b, Server, Tokens/s w tabeli Inference) 1.1.5.2. Oferowany System musi być wyposażony w akceleratory AI o potwierdzonym w benchmarku MLCommons wyniku dla modelu llama3.1-405b powyżej 100 Queries/s dla pojedynczego akceleratora (uzyskany wynik Systemu z kolumny llama3.1-405b, Server, Queries/s podzielony przez wartość z kolumny #of Accelerators i następnie podzielony przez wartość z kolumny #of Nodes) 1.1.5.3. Oferowany System musi być wyposażony w akceleratory AI o potwierdzonym w benchmarku MLCommons wyniku dla modelu mixtral-8x7b powyżej 15000 Tokens/s dla pojedynczego akceleratora (uzyskany wynik Systemu z kolumny mixtral-8x7b, Server, Tokens/s podzielony przez wartość z kolumny #of Accelerators i następnie podzielony przez wartość z kolumny #of Nodes)
2.	Wydajność	2.1. Łączna wydajność oferowanego Systemu musi gwarantować co najmniej 2 EFLOPS FP8 (2000 petaFLOPS FP8) trenowania modeli AI oraz co najmniej 4 EFLOPS FP8 (4000 petaFlops FP8) inferencji modeli AI. 2.2. Deklarowana wydajność oferowanego Systemu musi być potwierdzona przez producenta.
3.	Procesory	3.1. Oferowany System musi posiadać procesory w standardzie x86 - 64 bity do uruchamiania systemu operacyjnego. 3.2. Każdy węzeł przetwarzający zadania AI w ramach oferowanego Systemu musi być wyposażony w co najmniej 100 rdzeni procesorów x86 – 64 bity.
4.	Pamięć operacyjna	4.1. Oferowany System musi łącznie posiadać co najmniej 60 TB pamięci operacyjnej. 4.2. Każdy węzeł przetwarzający zadania AI w ramach oferowanego Systemu musi posiadać co najmniej 1 TB pamięci operacyjnej.
5.	Pamięć masowa	5.1. Oferowany System musi zawierać skalowalną pamięć masową dedykowaną do rozwiązań AI, certyfikowaną przez producenta GPU. 5.2. Pojemność oferowanego systemu pamięci masowej musi obejmować co najmniej 1 PB dostępnej surowej przestrzeni na dane. 5.3. Deklarowana wydajność systemu pamięci masowej musi zapewniać co najmniej 180 GB/S dla sekwencyjnego odczytu oraz 120 GB/s sekwencyjnego zapisu. 5.4. Oferowany system pamięci masowej musi umożliwiać instalację w standardowej szafie sprzętowej o rozmiarze do 4 RU. 5.5. Oferowany system pamięci masowej musi charakteryzować się nie większym zużyciem prądu niż 5 KW.

		5.6. Oferowany system pamięci masowej musi być wyposażony w co najmniej 16 interfejsów sieciowych o przepływności co najmniej 200 Gbit/s. 5.7. Oferowany system pamięci masowej musi wspierać standard RDMA oraz połączenia dedykowane dla GPU Systemu. 5.8. Oferowany system pamięci masowej musi umożliwiać konfigurację globalnych dysków Hot Spare. 5.9. Oferowany system pamięci masowej musi zawierać awaryjne zasilanie z akumulatora lustrzanej pamięci podręcznej o parametrach pojemnościowych i wydajnościowych zapewniających możliwość zapisu danych z pamięci podręcznej na przestrzeń dyskową w razie utraty zasilania i wymuszenia wyłączenia całego systemu. 5.10. Oferowany system pamięci masowej musi zapewniać redundantne, aktywne kontrolery. 5.11. Oferowany system pamięci masowej musi być wysoce dostępny i redundantny. W szczególności kontrolery, moduły/porty IO, kable, zasilacze, wentylatory itp. 5.12. Oferowany system pamięci masowej musi obsługiwać dyski samoszyfrujące (Self Encrypting Device) bez wpływu na wydajność. 5.13. Oprogramowanie systemu pamięci masowej musi zapewniać wysokowydajny, równoległy system plików wraz z kontrolą integralności danych.
6.	Dyski twarde	6.1. Każdy węzeł przetwarzający zadania AI oferowanego Systemu musi posiadać zainstalowane co najmniej 2 dyski min. 1TB SSD służące do uruchamiania systemu oraz co najmniej 8 dysków SSD NVMe o pojemności min. 3 TB.
7.	Wsparcie dla systemów operacyjnych	7.1. Oferowany System musi zostać dostarczony wraz systemem operacyjnym Linux, którego zaoferowana dystrybucja jest wspierana i certyfikowana przez producenta Systemu. 7.2. System operacyjny musi zawierać optymalizacje jądra systemu, w tym wsparcie dla bezpośredniego dostępu akceleratorów AI (tj. GPU) do pamięci masowych 7.3. Oferowany System musi umożliwić wdrożenie w pełni odizolowane od internetu (ang. Air-Gapped) w celu umożliwienia bezpiecznej aktualizacji systemu operacyjnego wraz z oprogramowaniem bazowym. Wymagane jest zapewnienie mechanizmów umożliwiających instalację i aktualizację oprogramowania bez konieczność zapewnienia dostępu do internetu czy usług w chmurze producenta Systemu.
8	Interfejsy sieciowe	8.1. Każdy węzeł przetwarzający zadania AI oferowanego Systemu musi posiadać co najmniej 4 interfejsy sieciowe Infiniband NDR o przepływności co najmniej 400 Gbit/s służące do komunikacji pomiędzy węzłami przetwarzającymi zadania AI.

		8.2. Każdy węzeł przetwarzający zadania AI musi posiadać 2 dedykowane interfejsy Infiniband NDR o przepływności co najmniej 200 Gbit/s do komunikacji z zewnętrzną pamięcią masową. 8.3. Każdy węzeł przetwarzający zadania AI musi posiadać min. 2 dedykowane interfejsy In-Band Management oraz jeden dedykowany interfejs Out-Band Management o przepływności co najmniej 200 Gbit/s do zarządzania. 8.4. Wykonawca dostarczy System z następującym ukończeniem urządzeń sieciowych wyprodukowanych przez producenta akceleratorów GPU: 8.4.1. Przełącznik InfiniBand, 64 porty NDR 400 Gb/s, 32 porty OSFP 400 Gb/s, łączna przepustowość min. 51 Tb/s, redundancja zasilania, zajętość szafy rack max 1U – 18 szt. 8.4.2. Unified Fabric Manager, zarządzanie siecią, 2 porty NDR 400 Gb/s, redundancja zasilania, zajętość szafy rack 1U – 4 szt. 8.4.3. Przełącznik sieci Ethernet 800 Gbe, 64 porty OSFP, 1 port SFP28, łączna przepustowość min. 51 Tb/s, redundancja zasilania, zajętość szafy rack max 2U – 4 szt. 8.4.4. Przełącznik sieci Ethernet 1GBase-T/100GbE, 48 portów RJ45 1Gb/s, 2 porty QSFP28 100Gb/s, przepustowość min. 440 Gb/s, redundancja zasilania, zajętość szafy rack max 1U – 7 szt. 8.4.5. Wszystkie porty w urządzeniach wymienionych powyżej muszą być aktywne bez konieczności zakupu dodatkowych licencji. 8.4.6. Komplet oryginalnych wkładek oraz interfejsów sieciowych zapewniających maksymalną przepustowość dla wszystkich urządzeń sieciowych. 8.4.7. Komplet kabli zapewniających realizację połączeń wewnętrznych pomiędzy węzłami Systemu.
9.	Akceleratory GPU	9.1. Oferowany System powinien umożliwiać wykorzystanie akceleratorów GPU najnowszej (na dzień składania ofert) generacji oferowanej przez producenta GPU będących częścią składową węzłów przetwarzania oferowanego Systemu. Oznacza to, że nawet jeżeli obecne wymagania wydajnościowe pozwolą na zastosowanie akceleratorów starszej generacji niż dostępne w dniu składania ofert, oferowany System musi umożliwić obsługę najnowszej generacji akceleratorów, która będzie dostępna w dniu składania ofert. Każdy węzeł Systemu musi zapewnić min. 1,4 TB pamięci GPU typu HBM3 łącznie.
10.	Gwarancja i wsparcie	10.1. Minimalnie 60 miesięczna gwarancja producenta Systemu w miejscu instalacji. Serwis będzie realizowany przez producenta Systemu lub jego partnera posiadającego ważną autoryzację do prowadzenia serwisu dla oferowanego Systemu. 10.2. Zgłoszenia przyjmowane w trybie 24x7x365 poprzez telefon oraz dedykowany kanał WWW. 10.3. Zgłoszenia incydentów powinny być klasyfikowane według ważności od 1-4:

		10.3.1. Poziom 1: Krytyczny incydent wstrzymujący możliwość realizacji kluczowych procesów biznesowych. Niedostępność krytycznych funkcjonalności bez możliwości zastosowania tymczasowego obejścia. Czas odpowiedzi dla najwyższego priorytetu to maksymalnie 1h. 10.3.2. Poziom 2: Incydent mający wpływ na realizację procesów biznesowych, istotne funkcje są niedostępne lub wymagają tymczasowego obejścia. Czas odpowiedzi to maksymalnie 2h. 10.3.3. Poziom 3: Incydent nie wstrzymujący realizacji kluczowych procesów biznesowych. Czas odpowiedzi to maksymalnie 4h. 10.3.4. Poziom 4: Incydent bez znacznego wpływu na realizację procesów biznesowych. Maksymalny czas odpowiedzi to 1 dzień roboczy. 10.4. Zagwarantowanie wymiany części zamiennych lub wsparcia w usuwaniu błędów, w przypadku awarii w miejscu instalacji: 10.4.1. W przypadku incydentu poziomu 1 do 4 godzin roboczych od zgłoszenia. 10.4.2. W przypadku incydentu poziomu 2 do końca następnego dnia roboczego od zgłoszenia. 10.4.3. W pozostałych przypadkach w ciągu 3 dni roboczych od zgłoszenia. 10.5. Zagwarantowanie możliwości zachowania przez Zamawiającego dysków pamięci masowej w przypadku konieczności wymiany dysków z powodu awarii w okresie trwania gwarancji. 10.6. Wykonawca zapewni dostęp do dedykowanego Zamawiającemu inżyniera wsparcia technicznego producenta Systemu. Inżynier wsparcia technicznego udzieli wsparcia z wykorzystaniem zdalnego dostępu do Systemu na każde żądanie, najpóźniej do 4 godzin roboczych od momentu zgłoszenia. 10.7. Zapewnienie instruktażu stanowiskowego w zakresie administrowania Systemem w siedzibie Zamawiającego lub miejscu instalacji. Odbiorcą instruktażu są administratorzy oraz personel odpowiadający za centrum przetwarzania danych. Instruktaż stanowiskowy musi obejmować wszystkie elementy administrowania Systemem, w szczególności infrastrukturą sieciową, obliczeniową, oprogramowaniem zarządczym oraz pamięcią masową. Ukończenie instruktażu powinno umożliwiać pracownikom Zamawiającego samodzielną administrację Systemem.

3.3. Wymagania dla dołączonego oprogramowania.

Lp.	Obszar	Wymaganie
1.	Bazowe oprogramowanie wsparcia AI	1.1. Bazowe oprogramowanie wsparcia AI musi posiadać wystawiony przez producenta tego oprogramowania lub organizację posiadającą stosowne prawa autorskie, certyfikat lub inny równoważny dokument poświadczający możliwość jego bezproblemowego i bezkonfliktowego wdrożenia i eksploatacji na infrastrukturze obliczeniowej oferowanego Systemu. W szczególności oferowane oprogramowanie musi być certyfikowane dla: 1.1.1. Architektury i generacji oferowanego Systemu. 1.1.2. Rodzaju i wersji oferowanego systemu operacyjnego. 1.1.3. Platformy orkiestracji kontenerów Kubernetes oraz kontenerów w technologii Docker lub OCI wdrażanych bezpośrednio na systemie operacyjnym, poza środowiskiem orkiestracji Kubernetes. 1.2. Licencja na oferowane oprogramowanie wsparcia AI musi zapewnić dostęp do kodówczy to to samo co wyżej? źródłowych oraz dokumentacji umożliwiającej budowanie aplikacji i ich rozwój z wykorzystaniem dostarczonych narzędzi i bibliotek. 1.3. Oferowane oprogramowanie AI musi posiadać funkcjonalności zarządzania infrastrukturą systemową, oprogramowaniem orkiestracyjnym, aby umożliwić skalowanie przetwarzania AI. W szczególności oprogramowanie musi posiadać: 1.3.1. Funkcjonalność zapewniającą automatyzację zarządzania sterownikami akceleratorów GPU udostępnioną jako Operator dla platformy orkiestracji kontenerów Kubernetes. 1.3.2. Katalog Kontenerów typu Docker wraz ze wsparciem automatyzacji wdrożenia w standardzie Helm Chart. 1.3.3. Funkcjonalność zapewniającą w klastrze Kubernetes dynamiczne przydzielanie akceleratorów GPU. Musi umożliwiać wielu użytkownikom i zadaniom współdzielenie pojedynczego GPU, alokując im ułamki mocy obliczeniowej w zależności od potrzeb. 1.3.4. Funkcjonalność zapewniającą w klastrze Kubernetes agregację akceleratorów GPU pozwalającą na łączenie mocy wielu GPU dla pojedynczego, intensywnego zadania (np. trenowanie modeli AI). 1.3.5. Funkcjonalność zapewniającą w klastrze Kubernetes automatyczne skalowanie poprzez dynamiczne dostosowanie alokacji zasobów GPU, CPU, pamięci w zależności od potrzeb uruchomionych zadań.

2.	Katalog Kontenerów	2. W Katalogu Kontenerów konieczne jest zapewnienie: 2.1. Serwera inferencji AI wspierającego MLfow, TensorFlow, PyTorch. 2.2. Narzędzia do gromadzenia i wizualizacji logów w ramach platformy. 2.3. Środowiska rozwoju modeli generatywnej sztucznej inteligencji. Środowisko musi wspierać rozwój multi modalnych modeli włączając Llama 2, Falcon, CLIP, LLAVA, GPT, T5, BERT, MoE, RETRO. Środowisko musi zapewniać pretrenowane, bazowe modele w obszarze automatycznego rozpoznawania mowy, przetwarzania języka naturalnego (NLP), zamiany tekstu na mowę oraz analizy obrazów (typu Computer Vision). Środowisko musi dostarczać funkcjonalności wspierające strojenie modeli typu Supervised Fine-Tuning (SFT) oraz Reinforcement Learning from Human Feedback (RLHF). 2.4. Środowisko budowy potoków analitycznych w oparciu o akceleratory GPU umożliwiające wysokopoziomowe programowanie w oparciu o język Python. Środowisko musi zawierać akcelerowane na GPU biblioteki cuDF, CuML, cuGraph, cuVS, RMM, RAFT lub równoważne.
3.	Katalog bazowych modeli	3. Oferowane oprogramowanie musi obejmować katalog bazowych modeli wraz z modelami dedykowanymi obrazowaniu medycznemu (przynajmniej 30 modeli) wykorzystywanych w praktyce w placówkach medycznych, w tym minimum po jednym modelu w następujących obszarach: 3.1. Interaktywny wstępnie wytrenowany model obrazowania medycznego 3D i 2D w zakresie segmentacji i anotacji anatomii człowieka wraz z możliwością strojenia modelu i dalszego trenowania na własnych zbiorach danych poprzez API. Wstępnie wytrenowany model sztucznej inteligencji to model uczenia głębokiego — wyrażenie algorytmu neuronowego podobnego do mózgu, który znajduje wzorce lub przewiduje na podstawie danych — który jest trenowany na dużych zestawach danych w celu wykonania określonego zadania. Może być używany w takiej postaci, w jakiej jest lub dalej dostrojony w celu dopasowania do konkretnych potrzeb aplikacji i/lub scenariuszy użycia. 3.2. Wstępnie wytrenowany model obrazowania medycznego w zakresie Tomografii Komputerowej służący do wolumetrycznej segmentacji 3d całego ciała człowieka na podstawie badań TK. 3.3. Wstępnie wytrenowany model obrazowania medycznego w zakresie Tomografii Komputerowej służący do segmentacji trzustki i guza trzustki z obrazu TK. 3.4. Wstępnie wytrenowany model do wolumetrycznego (3D) wykrywania guzków płucnych na obrazach Tomografii Komputerowej. 3.5. Wstępnie wytrenowany model do automatycznego wykrywania przerzutów w obrazach histopatologicznych.

		3.6. Wstępnie wytrenowany model do jednoczesnej segmentacji i klasyfikacji jąder w wielotkankowych obrazach histologicznych. 3.7. Wstępnie wytrenowany model do wolumetrycznej (3D) segmentacji podregionów guza mózgu z multimodalnych obrazów rezonansu magnetycznego.
4.	Środowisko uruchomieniowe medycznych modeli obrazowych	4.1 Oferowane oprogramowanie musi obejmować środowisko uruchomieniowe dla modeli medycznego obrazowania. 4.1. Biblioteki bazowe w języku Python z wsparciem dla środowiska pyTorch umożliwiające przetwarzanie obrazowania medycznego (w szczególności formatów DICOM, NIFTI, PNG, JPEG) 4.2. Narzędzie etykietowania obrazów w celu uczenia modeli AI, które umożliwia użytkownikowi tworzenie oznaczonych zbiorów danych 4.3. Środowisko budowania aplikacji obrazowania medycznego wykorzystujących modele AI. Środowisko musi być oparte o standard OCI - Open Container Initiative oraz wspierać co najmniej jeden otwarty standard pakietu aplikacyjnego obrazowania medycznego. 4.4. Środowisko uruchomieniowe aplikacji obrazowania medycznego budowanych z wykorzystaniem środowiska budowania aplikacji obrazowania medycznego (w pkt. 5.3). Środowisko uruchomieniowe musi dostarczać narzędzie CLI umożliwiające uruchamianie i testowanie wdrożonych aplikacji w sposób intuicyjny, niewymagający znajomości szczegółów implementacyjnych.