

Proof Copy ([1/5] w trakcie przeredagowywania)

Prof. dr inż. Jan Pająk

Komora Oscylacyjna

Monografia naukowa nr 2 z serii [1/5] "Zaawansowane urządzenia magnetyczne",
5-te wydanie, Wellington, Nowa Zelandia, 2007 rok,
ISBN 978-1-877458-02-6.

Copyright © 2007 by Prof. dr inż. Jan Pająk.

Wszystkie prawa zastrzeżone. Całość ani też żadna z części niniejszej monografii nie może zostać skopiowana, zreprodukowana, przesłana, lub upowszechniona w jakikolwiek sposób (np. komputerowy, elektroniczny, mechaniczny, fotograficzny, nagrania telewizyjnego, itp.) bez uprzedniego otrzymania wyrażonej na piśmie zgody autora lub zgody osoby legalnie upoważnionej do działania w imieniu autora. Od uzyskiwania takiej pisemnej zgody na kopiowanie tej monografii zwolnieni są tylko ci którzy zechcą wykonać jedną jej kopię wyłącznie dla użytku własnego nastawionego na podnoszenie wiedzy i dotrzymają warunków że wykonanej kopii nie użyją dla jakiegokolwiek działalności zawodowej czy przynoszącej dochód, a także że skopiowaniu poddadzą cały wybrany tom lub całą monografię - włącznie ze stroną tytułową, streszczeniem, spisem treści i rysunków, wszystkimi rozdziałami, tablicami, rysunkami i załącznikami.

Zarejestrowano w Bibliotece Narodowej Nowej Zelandii jako Depozyt Legalny. Wydano w Nowej Zelandii prywatnym nakładem autora w dwóch językach: angielskim i polskim.

Data ostatniej aktualizacji niniejszego tomu: 21 września 2007 roku. (W przypadku dostępu do kilku egzemplarzy/wersji tej monografii, rekomendowane jest zapoznanie się z egzemplarzem o najnowszej dacie aktualizacji!)

Niniejsza monografia stanowi raport naukowy z przebiegu badań autora. Stąd prezentacja wszelkich zawartych w niej materiałów posiadających wartość dowodową lub dokumentacyjną dokonana została według standardów przyjętych dla publikacji (raportów) naukowych. Szczególna uwaga autora skupiona była na wymogu odtwarzalności i możliwie najpełniejszego udokumentowania źródeł, tj. aby każdy naukowiec czy hobbysta pragnący zweryfikować lub pogłębić badania autora był w stanie dotrzeć do ich źródeł (jeśli nie noszą one poufnego charakteru), powtórzyć ich przebieg, oraz dojść do tych samych lub podobnych wyników.

Niniejsza monografia nr 2 należy do serii najważniejszych monografii naukowych autora. Może ona być czytana w odosobnieniu, lub czytana jako kolejny tom w całej owej serii. Cała seria tych monografii jest oznaczana symbolem [1/5] i nosi generalny tytuł "Zaawansowane urządzenia magnetyczne". Stanowi ona piąte wydanie owej serii. Jej poprzednim (czwartym) wydaniem była 18-tomowa monografia [1/4] opublikowana w 2001 roku i nosząca następujące dane edytorskie: "Zaawansowane urządzenia magnetyczne", Monografia, Wellington, Nowa Zelandia, 2001, ISBN 0-9583727-5-6, około 1800 stron tekstu (w tym 120 ilustracji i 7 tablic, w 18 tomach).

Wszelka korespondencja przeznaczona dla autora niniejszej monografii z serii [1/5] przed końcem 2007 roku powinna być kierowana na następujący jego adres w Korei:

Prof. Dr Eng. Jan Pająk,
Ajou University, South Korea

Natomiast od stycznia 2008 roku powinna ona być kierowana na następujący adres autora w Nowej Zelandii:

P.O. Box 33250, Petone 5046, New Zealand

Tel. domowy (2007 rok): +64 (4) 56-94-820; E-mail: jpajak@poczta.wp.pl lub janpajak@gmail.com

STRESZCZENIE tej monografii nr 2 z serii [1/5] "Zaawansowane urządzenia magnetyczne", ISBN 978-1-877458-04-0

Czym będzie "komora oscylacyjna"? Wyobraźmy sobie kryształową kostkę stanowiącą nowe urządzenie do produkcji super-silnego pola magnetycznego. Wyglądałaby ona jak idealnie ukształtowany kryształ jakiegoś przezroczystego minerału, lub jak sześcian wyszlifowany ze szkła i ukazujący swe wnętrze poprzez przezroczyste ścianki. Przy wielkości nie większej od poręcznej kostki Rubika wytwarzałaby ona pole setki tysięcy razy przewyższające pola produkowane dotychczas na Ziemi, włączając w to pola najsilniejszych współczesnych dźwigów magnetycznych czy najpotężniejszych elektromagnesów w laboratoriach naukowych. Gdybyśmy kostkę taką wzięli do ręki, wykazywałaby ona zdumiewające własności. Przykładowo mimo swych niewielkich rozmiarów byłaby ona niezwykle "ciężka" i przy jej przesterowaniu na pełny wydatek magnetyczny nawet najsilniejszy atleta nie byłby w stanie jej udźwignąć. Jej "ciężar" wynikałby z faktu, iż wytwarzane przez nią potężne pole magnetyczne powodowałoby jej przyciąganie w kierunku Ziemi i przez to do jej rzeczywistego ciężaru dodawałaby się wytworzona w ten sposób siła jej oddziaływań magnetycznych z polem ziemskim. Byłaby ona też oporna na nasze próby obracania i podobnie jak igła magnetyczna zawsze starałaby się zwrócić w tym samym kierunku. Gdybyśmy jednak zdołali ją obrócić w położenie dokładnie odwrotne do tego jakie sama starałaby się przyjmować, wtedy ku naszemu zdumieniu zaczęłaby nas unosić w powietrze. Sama jedna mogłaby więc napędzać nasze wehikuły.

Komora oscylacyjna posiada potencjał aby już wkrótce stać się jednym z najważniejszych urządzeń technicznych naszej cywilizacji. Jej zastosowania mogą być wszechstronne. Począwszy od akumulatorów energii o obecnie trudnej do wyobrażenia pojemności (np. komora o wielkości kostki do gry będzie w stanie zaspokoić zapotrzebowanie na energię ogromnych miast czy fabryk), poprzez urządzenia napędowe jakie umożliwią szybowanie w przestrzeni naszych wehikułów, osób, budynków a nawet mebli, a skończywszy na wypełnianiu funkcji prawie wszystkich naszych obecnych urządzeń przetwarzających energię, takich jak latarka, grzejnik, silnik spalinowy, ogniwo termoelektryczne, silnik elektryczny, generator elektryczności, transformator, magnes, oraz wiele innych. Znaczenie komory oscylacyjnej dla naszej sfery technicznej będzie mogło być tylko porównane do znaczenia komputerów dla naszej sfery intelektualnej.

Komora oscylacyjna wynaleziona została w pierwszych godzinach 3 stycznia 1984 roku. Natychmiast po opublikowaniu jej zasady działania, wielu hobbystów z kilku krajów świata rozpoczęło eksperymenty nad jej realizacją. Niestety, jest to urządzenie trudne do zbudowania, stąd hobbyści skompletowali jedynie model komory, ale nie byli w stanie otrzymać użytecznego prototypu. Jednakże ich bezspornym osiągnięciem było niezależne potwierdzenie iż zasada działania komory jest poprawna i daje się zrealizować na drodze technicznej.

Niezależnie od tych prób praktycznego zbudowania komory, podjąłem też badania porównawcze mające na celu teoretyczne potwierdzenie poprawności jej konceptu. Badania te zakończyły się całkowitym sukcesem wykazując poprawność idei komory zanim jeszcze program jej praktycznej realizacji na Ziemi zdążył zostać zainicjowany.

Niniejszy tom zestawia najważniejsze informacje o komorze oscylacyjnej. Omawia więc ona budowę, zasadę działania i postępy w praktycznej realizacji tego urządzenia. Dostarcza krótkiego przeglądu zastosowań komory, koncentrując się wszakże na zastosowaniach w urządzeniach napędowych (które będą stanowiły jedynie mały ułamek wszystkich jej zastosowań). Ujawnia też fakty jakie dokumentują iż urządzenie to zostało już zbudowane i czasami jest/było użytkowane na Ziemi. Tom ten reprezentuje więc opracowanie źródłowe dla wszystkich tych którzy zechcą zapoznać się z komorą oscylacyjną w celach badawczych, wynalazczych, czy po prostu aby poszerzyć swoje horyzonty.

Niezależnie od komory oscylacyjnej, tom ten omawia także dwa istotne jej zastosowania napędowe, tj. wehikuł czteropędnikowy oraz magnetyczny napęd osobisty. (Trzecie z najistotniejszych zastosowań napędowych tej komory, tj. magnokraft, omówione jest w rozdziale F, tom 3.)

Zanim jednak rozważania o komorze oscylacyjnej są zapoczątkowane, tom ten wyjaśnia także skąd wywodzi się nasza pewność, że na Ziemi nadchodzi właśnie czas zastosowań komory oscylacyjnej i używającej tą komorę magnetycznych urządzeń napędowych. Wyjaśnienie to zawarte jest w rozdziale B, który prezentuje tzw. "Tablicę Cykliczności", czyli odpowiednik słynnej Tablicy Mendelejewa opracowanej jednak dla urządzeń technicznych, a nie dla pierwiastków chemicznych. Owa Tablica Cykliczności wykazuje, że naturalnym następstwem dotychczasowego rozwoju napędów na Ziemi, jest właśnie obecna budowa magnokraftu i komory oscylacyjnej.

SPIS TREŚCI tej monografii 2 z serii [1/5] "Zaawansowane urządzenia magnetyczne", ISBN 978-1-877458-02-6.

Str: Rozdział:

- | | |
|---|--|
| 1 | Strona tytułowa |
| 2 | Streszczenie niniejszej monografii nr 2 |
| 3 | Spis treści niniejszej monografii nr 2
(odnotuj że spis treści całej serii monografii [1/5] zawarty jest w monografii nr 1) |

Monografia 2: Komora Oscylacyjna (ISBN 978-1-877458-02-6)

- | | |
|------|--|
| B-01 | B. PRAWO CYKLICZNOŚCI
TECHNICZNYM ODPOWIEDNIKIEM TABLICY MENDELEJEW |
| B-02 | B1. Trzy generacje magnokraftów |
| B-04 | B2. Podstawowy wymóg budowania sterowalnych napędów |
| B-04 | B3. "Trend omnibusa" a wygląd magnokraftów wszystkich trzech generacji |
| B-07 | Tablica B1 (Tablica Cykliczności) |
| C-08 | C. KOMORA OSCYLACYJNA |
| C-10 | C1. Dlaczego niezbędnym jest zastąpienie elektromagnesów przez komorę oscylacyjną |
| C-12 | C2. Historia komory oscylacyjnej |
| C-20 | C3. Zasada działania komory oscylacyjnej |
| C-21 | C3.1. Inercja elektryczna induktora stanowi siłę motoryczną dla oscylacji w tradycyjnym obwodzie oscylacyjnym z iskrownikiem |
| C-22 | C3.2. W zmodyfikowanym obwodzie oscylacyjnym z iskrownikiem inercji elektrycznej dostarczy induktancja iskry elektrycznej |
| C-24 | C3.3. Zestawienie razem dwóch zmodyfikowanych obwodów formuje komorę oscylacyjną wytwarzającą dipolarne pole magnetyczne |
| C-26 | C3.4. Igłowe elektrody |
| C-26 | C4. Przyszły wygląd komory oscylacyjnej |
| C-27 | C4.1. Trzy generacje komór oscylacyjnych |
| C-30 | C5. Matematyczny model komory oscylacyjnej |
| C-30 | C5.1. Oporność komory oscylacyjnej |
| C-30 | C5.2. Indukcyjność komory oscylacyjnej |
| C-31 | C5.3. Pojemność komory oscylacyjnej |
| C-32 | C5.4. Współczynnik motoryczny iskier i jego interpretacja |
| C-33 | C5.5. Warunek zaistnienia oscylacji we wnętrzu komory |
| C-33 | C5.6. Okres pulsowań pola komory oscylacyjnej |
| C-34 | C6. Jak komora oscylacyjna eliminuje wady elektromagnesów |
| C-35 | C6.1. Neutralizacja sił elektromagnetycznych |
| C-36 | C6.2. Niezależność wytwarzanego pola od ciągłości i efektywności dostawy energii |
| C-37 | C6.3. Eliminacja strat energii |
| C-38 | C6.3.1. Czy w komorze całe ciepło iskier będzie odżykiwalne |
| C-40 | C6.4. Spożytkowanie niszczyielskiego pola elektrycznego |
| C-41 | C6.5. Sterowanie amplifikujące okresu pulsowań pola |
| C-42 | C7. Dodatkowe zalety komory oscylacyjnej ponad elektromagnesami |
| C-42 | C7.1. Formowanie "kapsuły dwukomorowej" |

	zdolnej do sterowania swym wydatkiem magnetycznym bez zmiany ilości zawartej w niej energii
C-45	C7.1.1. Kapsuły dwukomorowe drugiej i trzeciej generacji
C-47	C7.1.2. Stopień upakowania komór oscylacyjnych i jego wpływ na wygląd kapsuł dwukomorowych i konfiguracji krzyżowych
C-48	C7.2. Formowanie "konfiguracji krzyżowej"
C-51	C7.2.1. Prototypowa konfiguracja krzyżowa pierwszej generacji
C-52	C7.2.2. Konfiguracje krzyżowe drugiej generacji
C-54	C7.2.3. Konfiguracje krzyżowe trzeciej generacji
C-55	C7.3. Nieprzyciąganie przedmiotów ferromagnetycznych
C-58	C7.4. Wielowymiarowa transformacja energii
C-58	C7.5. Nienawrotne oscylacje - unikalny atrybut komory umożliwiający akumulowanie przez nią nieograniczonych ilości energii
C-60	C7.6. Funkcjonowanie jako pojemny akumulator energii
C-60	C7.7. Prostota produkcji
C-61	C8. Wytyczne dla eksperymentów praktycznych nad komorą oscylacyjną
C-61	C8.1. Stanowisko badawcze
C-63	C8.2. Etapy, cele i metodyka budowy komory oscylacyjnej
C-70	C8.3. Przykłady tematów badawczych inicjujących prace nad komorą
C-72	C9. Przyszłe zastosowania komory oscylacyjnej
C-73	C10. Moje monografie poświęcone komorze oscylacyjnej
C-75	C11. Symbole, notacje i jednostki występujące w rozdziale C
C-76/89	Tablica C1 i 13 rysunków (C1 do C13)
CB-90	CB. PRZYSZŁE ZASTOSOWANIA KOMORY OSCYLACYJNEJ
CB-93	CB1. Przyszłe zastosowania komory oscylacyjnej jako akumulatora do bezspalinowych samochodów (czyli do tzw. „eco-cars”)
CB-93	CB2. Senator McCain obiecał nagrodę w wysokości 300 millions dolarów dla wynalazcy akumulatora energii który by wykazywał cechy komory oscylacyjnej
D-95	D. MAGNOKRAFT CZTEROPĘDNIKOWY
D-95	D1. Ogólna konstrukcja magnokraftu czteropędnikowego
D-97	D2. Działanie magnokraftu czteropędnikowego
D-98	D3. Własności magnokraftu czteropędnikowego
D-100	D4. Wygląd i wymiary magnokraftu czteropędnikowego
D-102	D5. Identyfikacja typu magnokraftu czteropędnikowego
D-103/105	Tablica D1 i rysunek D1
E-106	E. MAGNETYCZNY NAPĘD OSOBISTY
E-107	E1. Standardowy kombinezon napędu osobistego
E-109	E2. Działanie magnetycznego napędu osobistego
E-111	E3. Kombinezon z pędnikami głównymi w naramiennikach
E-112	E4. Wersja napędu osobistego z poduszkami wokół bioder
E-112	E5. Osiągi napędu osobistego
E-113	E6. Podsumowanie atrybutów magnetycznego napędu osobistego
E-116 /119	4 rysunki (E1 do E4)

Uwagi:

(1) Niniejsza monografia jest kolejną publikacją z całej serii 18 monografii naukowych autora. Każdy rozdział i podrozdział z owej serii oznaczany jest następną literą

alfabetu. Rozdziały oznaczane innymi literami niż te podane w powyższym spisie treści znajdują się w odrębnych monografiach (tomach) tej serii. Pełny spis treści wszystkich 18 monografii (tomów) tej serii przytoczony jest w pierwszej monografii (tome 1).

(2) Istnieje też angielskojęzyczna wersja niniejszej monografii. W przypadku więc trudności ze zdobyciem jej polskojęzycznej wersji, czytelnicy znający dobrze język angielski mogą się z nią zapoznać w języku angielskim. Mogą też ją polecać do poczytania swoim znajomym nie znającym języka polskiego.

(3) Obie wersje tej monografii [1/5], tj. angielskojęzyczna i polskojęzyczna, używają tych samych ilustracji. Jedynie podpisy pod ilustracjami wykonane są w odmiennym języku. Dlatego w przypadku trudności z dostępem do polskojęzycznych ilustracji, przeglądając można angielskojęzyczne ilustracje. Warto również wiedzieć, że powiększenia wszystkich ilustracji z monografii [1/5] są dostępne w internecie. Dlatego aby np. przeanalizować powiększenia tych ilustracji, warto je przeglądać bezpośrednio w internecie. Dla ich znalezienia czytelnik powinien odszukać dowolną totaliztyczną stronę którą ja autoryzuję, np. poprzez wpisanie słowa kluczowego **"totalizm"** do dowolnej wyszukiwarki (np. do www.google.com), potem zaś - kiedy owa totaliztyczna strona się ukaże, czytelnik powinien albo uruchomić jeszcze jedną stronę nazywającą się "tekst_1_5.htm" (z egzemplarzami i ilustracjami monografii serii [1/5]) dostępną na tym samym serwerze, albo też wybrać opcję "tekst [1/5]" z menu owej totaliztycznej strony. Proszę też odnotować że wszystkie totaliztyczne strony pozwalają na załadowanie do swego komputera darmowych egzemplarzy całej niniejszej serii monografii [1/5].

(4) W przypadku wykonywania wydruku niniejszej monografii [1/5] z wersji internetowej, lub z wersji wordprocessorowej dostarczonej na dyskietce, podana w tym spisie treści numeracja stron niekoniecznie musi się pokrywać z numeracją stron na wydruku i dlatego może wymagać wprowadzenia korekcji. Wynika to z faktu że formatowanie niniejszego tomu tej monografii miało miejsce dla fontu "Arial" (12 punktów) oraz dla standardowych nastaw wordprocessora "Word XP" pracującego pod WINDOWS XP. Dla innych fontów, oraz dla innych wersji Word'a, formatowanie to niekoniecznie musi dawać takie same zagospodarowanie stron. Warto tutaj też pamiętać, że sprowadzone z internetu wersje źródłowe tej monografii, czytelnicy mogą sobie sami sformatować na dowolny font i dowolną gęstość, tak aby możliwie najlepiej dostosowane one były do ich nawyków czytelniczych i do ostrości ich wzroku.

(5) Aktualizacja i przeredagowanie niniejszego piątego wydania [1/5] tej monografii będzie postępowo. Czytelnik może się zorientować ze spisu treści które jej podrozdziały zostały już przeredagowane, lub właśnie są w trakcie przeredagowywania, bowiem na spisie treści są one oznaczone przez dodanie im symbolu # przed ich numerem. Pozostałe podrozdziały tej monografii ciągle należy poznawać w brzmieniu w jakim zostały one sformułowane dla poprzedniego wydania [1/4] tej monografii.

(6) W celu udoskonalenia struktury tej serii monografii [1/5], kolejność jej rozdziałów w niektórych monografiach została nieco zmieniona w stosunku do tej kolejności występującej w monografii [1/4]. I tak w monografiach 10 i 11 rozdziały K, L, M i N z [1/4], w tej serii monografii [1/5] są oznaczane odpowiednio K, M N i L.

PRAWO CYKLICZNOŚCI TECHNICZNYM ODPOWIEDNIKIEM TABLICY MENDELEJEW

Motto: "Przyglądnij się przeszłości a zobaczysz przyszłość."

W artykule [1A4] opublikowane zostało ogromnie ważne odkrycie, jakie później nazwałem "Prawem Cykliczności" w rozwoju napędów. Prawo Cykliczności wnosi taki sam porządek i symetryczność do rozwoju napędów, jak Tablica Mendelejewa wniosła do naszego poznania budowy materii. Prawo to stwierdza, że "budowanie urządzeń napędowych podlega generalnym regułom symetrii (cykliczności), tak że znając działanie napędów odkrytych w przeszłości możliwym się staje przewidywanie działania nowych napędów jakich zbudowanie nastąpi dopiero w przyszłości". Rozszerzony opis Prawa Cykliczności opublikowany został w monografiach [1a], [3] i [6]. **Tablica B1** w niniejszej monografii ilustruje działanie tego prawa.

Odkrycie Prawa Cykliczności nastąpiło gdy uświadomiłem sobie, że wszelkie napędy można podzielić na dwie drastycznie różne klasy, które zostały nazwane "silnikami" i "pędnikami". **Silniki** wytwarzają tylko ruch względny jednych części danej maszyny, względem jej innych części. Ich przykładem może być: silnik w pralce (jaki obraca bęben względem obudowy) czy silnik w tokarce. Silniki nie są w stanie wytworzyć absolutnego ruchu całych obiektów względem otoczenia, chociaż często dostarczają one mechanicznej energii dla urządzeń wytwarzających taki ruch absolutny (np. w samochodzie - koła, a nie silnik, powodują jego jazdę po gruncie, chociaż to silnik dostarcza kołom niezbędnej energii mechanicznej). Zupełnie odmienne od silników są **pędniki**, które powodują ruch absolutny całych obiektów w otaczającym je ośrodku naturalnym. Przykładem pędników mogą być: koło samochodowe, śmigło lotnicze, czy dysza rakietowa. Zauważ że pędniki zawsze są w stanie działać w naturalnym ośrodku i stąd należy je wyraźnie odróżniać od "silników liniowych" w których część wspierająca ruch (np. szyna, prowadnica, itp.) została wydłużona na odpowiednią odległość. Przykładowo napęd kolei żelaznych, kolei magnetycznych, kolei linowych, wind, czy przenośników taśmowych, z definicji należy zaliczać do silników liniowych, nie zaś do pędników.

Po dokonaniu podziału napędów na silniki i pędniki spostrzegłem, że dla każdego rodzaju czynnika roboczego zawsze budowana jest para bliźniaczych urządzeń napędowych, z których pierwsze jest silnikiem, a drugie - pędnikiem (patrz tablica B1). Oba takie bliźniacze urządzenia działają na niemalże identycznej zasadzie, zaś pierwsze ich wersje konstrukcyjne są także ogromnie podobne. Przykładem owych bliźniaczych par mogą być: wiatrak i żagiel, czy silnik spalinowy i rakiet (tj. jeśli usunie się tłok z cylindra silnika spalinowego wtedy uzyska się dyszę rakiety). Analiza dat zbudowania poszczególnych urządzeń takich par wskazuje też, że odstęp czasowy pomiędzy nimi z reguły nie przekracza 200 lat. Na powyższym spostrzeżeniu oparte zostało uproszczone sformułowanie Prawa Cykliczności które mówi, że **"dla każdego silnika w przeciągu do 200 lat budowany jest odpowiadający mu pędnik"**. Sformułowanie to postuluje, że jeśli istnieje jakiś silnik, który dotychczas nie posiada bliźniaczego pędnika, odpowiadający mu pędnik powinien zostać zbudowany nie później niż około 200 lat od opracowania owego silnika. Wszyscy znamy taki odosobniony silnik. Jest nim zwyczajny silnik elektryczny, zbudowany przez Jacobi'ego około 1836 roku, w którym ruch wytwarzany zostaje przez siły przyciągających i odpychających oddziaływań magnetycznych. Jeśli więc Prawo Cykliczności działa, jeszcze przed rokiem

2036 silnik elektryczny powinien doczekać się zbudowania swego następcy w postaci pędnika magnetycznego zdolnego do napędzania magnokraftów opisanych w rozdziale F.

Po wykryciu istnienia bliźniaczych par silnik-pędnik uporządkowałem je w formę tzw. "Tablicy Cykliczności", w niniejszej monografii pokazanej jako tablica B1. W tablicy tej dwie pary silnik-pędnik formują jedną generację napędów eksploatujących daną właściwość czynnika roboczego, zaś trzy kolejne takie generacje zamykają pełny cykl wykorzystania danego czynnika roboczego i wyczerpują listę urządzeń napędowych jakie można zbudować w oparciu o ten czynnik. Omawiana tablica ma tę właściwość, że pomiędzy wszystkimi jej elementami występuje wyraźnie dająca się zaobserwować symetryczność. Symetryczność ta stanowi esencję Prawa Cykliczności, umożliwia ona bowiem przenoszenie (ekstrapolację) cech pomiędzy różnorodnymi urządzeniami napędowymi (podobnie jak to można czynić z pierwiastkami w Tablicy Mendelejewa). Dzięki niej, cechy napędów których zbudowanie nastąpi dopiero w przyszłości mogą zostać ekstrapolowane z cech napędów już istniejących.

Po odkryciu i pełnym rozpracowaniu Prawa Cykliczności, podążyłem za zawartymi w nim regułami symetryczności, i w ten sposób precyzyjnie rozpracowałem szczegóły budowy i działania magnokraftu. Stąd opisy z niniejszej monografii, jak również sformułowania wszystkich innych moich publikacji poświęconych zaawansowanym napędom magnetycznym, stanowią jedynie konsekwencję praktycznego wykorzystania wskazówek wynikających z Prawa Cykliczności.

Pełny opis Prawa Cykliczności oraz zasad budowy Tablic Cykliczności jest obszernym tematem jakiego wyczerpujące zaprezentowanie wymagałoby napisania oddzielnej monografii o objętości zbliżonej do niniejszego opracowania. W monografii [1a] tylko przedstawienie najistotniejszych aspektów tego prawa zajęło rozdział o objętości 34 stron. Z uwagi więc na konieczność ograniczania objętości niniejszej monografii, omówione tu zostaną jedynie zagadnienia albo posiadające bezpośredni związek z jej główną tezą, albo też pomocne w zrozumieniu przytoczonych później rozważań. Czytelnikom pragnącym zapoznać się też z działaniem Prawa Cykliczności dla urządzeń energetycznych, rekomendowane jest przeczytanie podrozdziału K1, a także następnego (trzeciego) wydania monografii [6] lub angielskojęzycznej monografii [1a] (w przyszłości planowane jest też opublikowanie następnego wydania [1/5] niniejszej polskojęzycznej monografii [1/4] zawierające już poszerzone ujęcie Prawa Cykliczności).

B1. Trzy generacje magnokraftów

W każdym urządzeniu napędowym obecne muszą być trzy następujące składniki: (1) czynnik roboczy, (2) wymiennik energii, oraz (3) przestrzeń robocza.

Czynnik roboczy jest to "medium" wykorzystywane w danym napędzie, którego funkcją jest absorbowanie jednego rodzaju energii i późniejsze wyzwolenie tej energii w formie oddziaływań siłowych zdolnych do wytworzenia ruchu. Przykładami czynników roboczych są: siły mechanicznej sprężystości (w łuku lub katapulcie), przepływająca woda (w kole młyńskim), para wodna (w silniku parowym), gazy spalinowe (w rakiecie kosmicznej) czy pole magnetyczne (w silniku elektrycznym).

Wymiennik energii jest to przestrzeń lub urządzenie w danym napędzie, w którym następuje wytworzenie czynnika roboczego i w którym czynnik ten absorbuje swoją energię początkową jaka następnie zostanie przez niego wyzwolona w celu wytworzenia określonego typu ruchu. Przykładami wymiennika energii mogą być: kocioł parowy ze silnika parowego, lub cewki zwojów elektromagnesu ze silnika elektrycznego.

Przestrzeń robocza jest to przestrzeń lub urządzenie w danym napędzie, gdzie następuje właściwe wytworzenie ruchu. W przestrzeni tej energia zawarta w czynniku roboczym zamieniona zostaje na pracę dostarczenia ruchu do napędzanego obiektu. Przykładami przestrzeni roboczej mogą być: przestrzeń pomiędzy tłokiem i cylindrem w silniku

parowym, dysza rakiety kosmicznej, szczeliny między łopatkami w turbinie parowej, czy szczelina pomiędzy wirnikiem i stojanem w silniku elektrycznym.

Z analiz urządzeń napędowych dotychczas skompletowanych na Ziemi wynika, iż jedynie trzy rodzaje "medium" nadają się do użycia jako czynnik roboczy. Są to: (1) obieg oddziaływań siłowych, (2) obieg masy, oraz (3) obieg linii sił pola magnetycznego. Stąd wszystkie dotychczas istniejące czynniki robocze mogą zostać zakwalifikowane do jednej z tych trzech generalnych kategorii (patrz pierwsza kolumna w tabeli B1), zależnie od tego które z powyższych trzech obiegów one reprezentują. Ponieważ podczas rozwoju naszej cywilizacji, owe trzy kolejne kategorie czynników roboczych były odkrywane we wyszczególnionej powyżej kolejności, możemy więc mówić o trzech erach w historii naszej myśli technicznej, w każdej z których jeden z powyższych czynników roboczych był szczególnie dominujący. I tak w starożytności i średniowieczu panowała era czynników roboczych bazujących na obiegu siły (np. siły bezwładności i reakcji w kole zamachowym, siły sprężystości w sprężynie zegarowej). Od czasu wynalezienia maszyny parowej (1769 rok) aż do dzisiaj, dominowała era czynników roboczych bazujących na obiegu masy (np. powietrze dla śmigła lotniczego, woda dla śruby okrętowej, gazy spalinowe dla pędnika odrzutowca). Obecnie zaczynamy jednak zbliżać się do trzeciej ery, w której dominującym zaczyna być obieg linii sił pola magnetycznego. Do chwili obecnej zbudowaliśmy tylko jeden i to najbardziej prymitywny reprezentant tych napędów przyszłości, tj. silnik elektryczny, który wykorzystuje obieg wytwarzanych przez siebie pól magnetycznych. Niemniej już wkrótce cały szereg bardziej zaawansowanych napędów tego typu zostanie urzeczywistnionych (ich opisy zawarte są w rozdziałach D, E i F tej monografii). Natomiast w dalszej przyszłości zbudowane będą nawet jeszcze bardziej zaawansowane ich wersje, jakie opisane są w rozdziałach L i M.

Dla każdego typu czynnika roboczego, aż trzy różne generacje urządzeń napędowych mogą zostać zbudowane - patrz tabela B1. W każdej z nich kolejne atrybuty czynnika roboczego zostają wykorzystane jako nośniki energii. Pierwsza z tych generacji zawsze wykorzystuje jedynie **oddziaływania siłowe** (np. popychanie, pociąganie, ciśnienie, ssanie, odpychanie, przyciąganie) wytwarzane przez czynnik roboczy. W drugiej generacji, na dodatek do tych oddziaływań siłowych, również **inercyjność** czynnika roboczego jest wykorzystywana (np. uderzenie, odrzut). Trzecia generacja urządzeń napędowych działających na danym czynniku roboczym wykorzystuje nie tylko jego oddziaływania siłowe i inercyjność, ale także wytwarzany przez niego **impakt energii wewnętrznej** (np. spężystości czy ciepła).

Z tabeli B1 wynika jednoznacznie iż silnik elektryczny i magnokraft reprezentują jedynie pierwszą i najbardziej prymitywną generację napędów bazujących na obiegu linii sił pola magnetycznego. Jedynym bowiem atrybutem pola jaki napędy te wykorzystują to siła odpychania lub siła przyciągania magnetycznego. Stąd po skompletowaniu tej pierwszej generacji magnokraftów, nasza cywilizacja przystąpi do kompletowania ich drugiej i trzeciej generacji. W każdej z nich, niezależnie od magnokraftu, aż cztery odmienne urządzenia należące do dwóch kolejnych par silnik-pędnik mogą zostać zbudowane (patrz górna część tabeli B1). Działanie owych zaawansowanych urządzeń napędowych będzie nie tylko wykorzystywało przyciągające i odpychające oddziaływania magnetyczne, ale także takie złożone zjawiska jak technicznie zaindukowaną telekinezę (jaka wyzwalana jest poprzez magnetyczny odpowiednik inercji - patrz objaśnienia w podrozdziałach H6.1 i L1, oraz w monografiach [1a], [3], [5] i [6]) oraz zmiany w upływie czasu (czas z kolei jest odpowiednikiem energii wewnętrznej pola magnetycznego - patrz podrozdziały H9.1 i M1 oraz opisy w monografii [8]). Stąd razem z magnokraftem pierwszej generacji, w toku rozwoju technicznego, nasza cywilizacja dorobi się aż trzech różnych generacji tych wehikułów, napęd drugiej i trzeciej generacji których wywoływał będzie zjawiska telekinezy oraz zjawiska zmian w upływie czasu. Aby wyrazić powyższe w terminologii z niniejszej monografii: druga i trzecia generacja magnokraftów zdolna będzie do działania w, odpowiednio, konwencji telekinetycznej i konwencji podróży w czasie. Następny podrozdział wyjaśni bliżej co należy rozumieć przez owo sformułowanie.

B2. Podstawowy wymóg budowania sterowalnych napędów

Jedną z podstawowych zasad fizyki, nazywaną "Zasadą Zachowania Pędu", stwierdza że gdziekolwiek jakiś system mas poddany zostaje działaniu wyłącznie sił wewnętrznych, jakie jedne masy tego systemu wywierają na inne, wtedy całkowity wektor pędu owego systemu pozostaje bez zmiany. Konsekwencją odniesienia tej zasady do napędów jest, że czynnik roboczy zawsze musi być w nich cyrkulowany przez otoczenie (w pednikach) lub przez element (w silnikach) w stosunku do którego ruch powinien zostać wytworzony. Powyższe reprezentuje "podstawowy wymóg cyrkulowania czynnika roboczego poprzez otoczenie w celu otrzymania sterowalnego napędu". Wymóg ten jest spełniony we wszystkich użytecznych komercyjnie napędach jakie ludzie wytworzyli dotychczas, nawet jeśli czasami przyjmuje on formę niebezpośrednią (np. w rakietach kosmicznych gdzie paliwo jest najpierw pobierane z otoczenia i umiejscawiane w zbiornikach rakiety, a następnie podczas lotu jest spalane i wyrzucane {tj. cyrkulowane} z powrotem do otoczenia).

Czasami konstruktor czy wynalazca jakiegoś urządzenia napędowego zignoruje wymóg aby czynnik roboczy cyrkulowany w nim został poprzez otoczenie. W takim przypadku wytwarzany przez ten napęd ruch jest niesterowalny i stąd nie może być stosowany do dla potrzeb wymagających precyzji i celowości. Urządzenie wytwarzające taki niesterowalny ruch nazywane tutaj będzie semi-napędem (tj. semi-silnikiem lub semi-pędnikiem). Semi-napędy stosunkowo łatwo przekonstruowane mogą być w napędy. Jedyne co w tym celu konieczne, to zapewnić cyrkulowanie ich czynnika roboczego poprzez otoczenie. Jako ich przykład można przytoczyć spadachron, który w starej "okrągłej" wersji nie cyrkulującej powietrza jest semi-pędnikiem. Jeśli jednak zorganizować cyrkulowanie w nim czynnika roboczego, poprzez uformowanie spadochronu na kształt lotni lub skrzydła, wtedy otrzymuje się sterowalny spadochron który umożliwia ślizganie się w powietrzu do dowolnego wybranego przez skoczka punktu docelowego. Innym przykładem semi-pędnika ciągle oczekującego takiej modyfikacji jest balon na gorące powietrze. Jeśli balon taki zaopatrzyć w sterowalną dyszę umiejscowioną z tyłu w bocznej części jego powłoki (dla lepszej aerodynamiczności powłokę tą dobrze byłoby budować we formie jakby wydłużonego sterowca), ten najstarszy ludzki wehikuł latający z semi-pędnika zamieni się w pędnik i będzie mógł latać w zamierzonym kierunku poprzez proste działania sterujące jego pilota zmieniające kierunek i wydatek strumienia powietrza wylatującego przez tą dyszę. Taka niewielka modyfikacja może zmienić balony na gorące powietrze w najbardziej proste, tanie, przyjemne, a jednocześnie efektywne środki transportu osobistego, znacznie wygodniejsze od motocykli i samochodów.

Zamiana semi-pędnika w pędnik z reguły nie wymaga jakiegokolwiek zasadniczej zmiany w jego konstrukcji. Z tego też powodu, w świetle Prawa Cykliczności zakładali będziemy, że dany rodzaj napędu został skompletowany bez względu na to czy jego końcowa czy semi-koncowa forma została dotychczas zbudowana.

B3. "Trend omnibusa" a wygląd magnokraftów wszystkich trzech generacji

W potocznym języku jedno ze znaczeń terminu "konwencja" brzmi "jednoznacznie zdefiniowane zachowanie". W niniejszej monografii ów termin zostanie więc przyjęty dla opisu ściśle zdefiniowanego zachowania się wehikułu latającego. Stąd począwszy od tego miejsca przez wyrażenie "konwencja działania wehikułu" rozumieli będziemy nazwę głównej zasady wykorzystanej w danym momencie przez określony wehikuł dla spowodowania swego lotu. Aby lepiej zrozumieć potrzebę wprowadzenia tego terminu, posłużmy się przykładem nie tak dawno opracowanego na Ziemi promu kosmicznego (np. "Columbia"). Prom taki zaprojektowany został z możliwością lotów na trzech zasadach, tj.: (1) jako rakieta, (2) jako szybowiec, lub (3) jako bezwładny satelita ziemski. Aby więc precyzyjnie określić którą z tych

trzech zasad w danym momencie określony prom wykorzystuje, koniecznym jest właśnie użycie pojęcia konwencji (np. sprecyzowanie "prom ten właśnie leci w konwencji szybowca").

W przypadku wehikułów jakich zbudowanie nastąpi dopiero w przyszłości, ich konwencja lotu posiada ogromne znaczenie. Wynika to z generalnego trendu w rozwoju urządzeń napędowych jaki w tym opracowaniu nazwany zostanie "trendem omnibusa". Aby lepiej opisać czym jest ów trend posłużymy się przykładem hipotetycznego samolotu jaki nazwiemy tu "omnibusem". Omnibus przyjmie kształt tuby z otwartymi końcami. Kształt ten umożliwia mu więc połączenie w pojedynczym wehikule działania aż trzech różnych generacji pędników wykorzystujących obieg masy, tj. szybowca, poduszkowca i odrzutowca. Kiedy więc omnibus leci na dużych wysokościach może on wygasić spalanie paliwa i szybować w powietrzu jak pędniki pierwszej generacji z obiegiem masy (tj. szybowce). W tablicy B1 pędniki te reprezentowane są przez żagiel. Kiedy omnibus skieruje strumień swoich gazów odrzutowych ku ziemi, zaczyna działać jak poduszkowiec ślizgając się poziomo tuż przy powierzchni gruntu. Podczas takiego działania reprezentuje on więc drugą generację pędników z obiegiem masy. Omnibus może też działać jako odrzutowiec, przecinając powietrze swoim tubiastym korpusem wyrzucającym z tyłu gazy odrzutowe. W tym więc przypadku reprezentuje on trzecią generację pędników z obiegiem masy.

Powyższe ujawnia, iż aby precyzyjnie określić na jakiej zasadzie omnibus działa w określonym momencie czasu, wprowadzenie konwencji staje się konieczne. Po jego wprowadzeniu, z łatwością będziemy mogli opisać zachowanie tego wehikułu mówiąc po prostu iż leci on albo w konwencji szybowca, poduszkowca, albo też odrzutowca. W każdej z nich ten sam omnibus zachowuje się jak całkowicie odmienna generacja napędów z obiegiem masy.

Dotychczas zgromadzone doświadczenia w zakresie użytkowania różnych wehikułów latających wykazują, iż wszystkie trzy kolejne generacje napędów z obiegiem masy nawzajem się uzupełniają. Stąd też obecne pędniki trzeciej generacji, takie jak przykładowo rakiety, nie tylko że nie są w stanie zastąpić pędników pierwszej i drugiej generacji, takich jak szybowiec czy poduszkowiec, ale nawet wprowadzają one zwiększone zapotrzebowanie na równoczesne wykorzystywanie tych niższych od siebie pędników. Jednym z przykładów takiego zapotrzebowania jest już wspomniany poprzednio prom kosmiczny (Columbia) który musi działać zarówno jako rakietę jak i jako szybowiec. Z drugiej strony, nasza rosnąca wiedza na temat systemów napędowych dostarcza nam już obecnie coraz większych możliwości technicznych aby nowobudowane wehikuly zaopatrywać w urządzenia napędowe umożliwiające równoczesne używanie kilku różnych konwencji lotów. Przykładem mogą tu być współczesne samoloty wojskowe, które zaopatrywane są w możliwości lotów odrzutowych, plus równoczesne możliwości pionowego startu (tj. działania jako poduszkowce), oraz równoczesne możliwości szybowania. Powyższe uświadamia iż budowanie omnibusów jest naturalnym trendem technicznym jaki z czasem będzie się tylko pogłębiał.

Powyższe można wyrazić w formie generalnej zasady o następującej treści: **"We wysoko rozwiniętych cywilizacjach trend omnibusa staje się tak dominujący, iż budowanie wyższych generacji wehikułów latających osiągnęte zostaje poprzez dodawanie dalszych konwencji lotu do już istniejących wehikułów niższej generacji wykorzystujących ten sam czynnik roboczy."**

Najsilniejszy wpływ na naszą cywilizację trend omnibusa wywrze gdy opracowana zostanie druga i trzecia generacja wehikułów z napędem magnetycznym. Owe dwa wysoko zaawansowane wehikuly nie będą bowiem budowane jako zupełnie nowe i całkowicie odmienne statki, ale raczej jako dodatkowo usprawnione wersje zwykłego magnokraftu. Ich kształt, wygląd zewnętrzny, zagospodarowanie przestrzeni wewnętrznej, a także jedna z konwencji działania (tj. konwencja lotów magnetycznych) będą identyczne do tych ze zwykłego magnokraftu. Jedyną różnicą jaką te zaawansowane wehikuly będą wykazywały w stosunku do zwykłego magnokraftu, to iż niezależnie od lotów w konwencji magnetycznej, gdy zajdzie potrzeba będą one zdolne do lotów w konwencji telekinetycznej (wehikuly drugiej generacji), lub do lotów w konwencji telekinetycznej albo konwencji podróży w czasie (wehikuly trzeciej

generacji). Aby więc podkreślić że oba te wysoko zaawansowane wehikuly zrodziły się ze zwykłego magnokraftu i mogą ciągle latać w konwencji zwykłego magnokraftu, w niniejszej monografii będą one nazywane magnokraftem drugiej generacji (tj. wehikuł zdolny do lotów w konwencji magnetycznej lub konwencji telekinetycznej) oraz magnokraftem trzeciej generacji (tj. wehikuł zdolny do lotów w konwencji magnetycznej, konwencji telekinetycznej, lub konwencji podróży w czasie). Dla przeciwstawienia się tym wehikulom, zwykły magnokraft (tj. ten opisany w podrozdziale A2 i w rozdziale F), jaki zdolny będzie jedynie do lotów w konwencji magnetycznej, będzie tu nazywany magnokraftem pierwszej generacji lub po prostu magnokraftem.

W tym miejscu powinno zostać podkreślone, iż każdy z tych wehikulów wyższej generacji w danej chwili może używać tylko jednej konwencji lotów. Dla przykładu jeśli magnokraft drugiej generacji leci w konwencji magnetycznej, jego zdolności telekinetyczne muszą zostać wyłączone. Gdy jednak włączy on swój napęd telekinetyczny, wtedy równocześnie jego siły przyciągania i odpychania magnetycznego muszą zostać wygaszone.

Kierunek upływu czasu (kierunek udoskonaleń napędów spowodowany upływem czasu)									
3.	Napędy z obiegiem pola magnetyczn.	3.	Energia wewn.	?	wehikuł czasu:2300	?	?	Przy- łość V	
		2.	Inercja pola	silnik teleki:2036	magnokr.telek:2200	?	?		
		1.	Siła pola	silnik elektr:1836	magnokraft:ok.2036	silnik pulsarowy	statek gwiazdzisty		
2.	Napędy z obiegiem masy	3.	Energia wewn.	silnik parowy:1769	odrzutowiec: 1939	sil. spalinow:1867	rakieta: 1942	Teraź- niej- szość V	
		2.	Inercja masy	silnik pneuma:1860	poduszkowiec: 1959	masz. atmosf:1712	śmigło: 1903		
		1.	Siła ciśnienia	wiatrak: 1191	żagiel: około 1390	puszka Vidi: 1860	balon: 1863		
1.	Napędy z obiegiem siły mechaniczn.	3.	Sprężystość	wiertło inercyjne	katapulta	sprężyna	piłka	Postęp	
		2.	Inercja	koło garncarskie	taran bitewny	koło zamchowe	proca		
		1.	Odział. siłowe	korba napędowa	tyczka flisarska	kierat	koło		
E	Rodzaj	Ge	Nośnik	Napędy	silniki 1 pary	pedniki 1 pary	silniki 2 pary	pedniki 2 pary	Postęp
r	czynnika	ne	energii		(ruch względny)	(ruch absolutny)	(ruch względny)	(ruch absolutny)	
a	roboczego	ra							
		c	Rozwiązania	Pierwsza para silnik-pędnik (przestrz. robocza oddzielona od wytwornika)		Druga para silnik-pędnik (przestrzeń robocza w wytworniku czynnika robocz.)			
		ja	techniczne						

Tablica B1. Tablica Cykliczności sporządzana dla napędów ziemskich. Stanowi ona rodzaj "Tablicy Mendelejewa", tyle tylko że obowiązującej dla urządzeń napędowych zamiast dla pierwiastków chemicznych. Jej sformułowanie ujawnia że budowa kolejnych napędów ziemskich podlega prawom generalnej (DeBroglie'wskiej) symetrii, których działanie zezwala na przenoszenie (ekstrapolację) istotnych cech pomiędzy poszczególnymi urządzeniami. To z kolei umożliwia przewidzenie zasad działania, cech oraz przybliżonych dat uruchamiania urządzeń napędowych dotychczas jeszcze nie zbudowanych na Ziemi. Tablica ta powstała poprzez odłożenie na jej pionowej osi wszystkich czynników roboczych wykorzystywanych w działaniu poszczególnych generacji urządzeń napędowych, zaś na jej poziomej osi wszystkich możliwych urządzeń napędowych budowanych dla danego czynnika roboczego. Jej pole robocze przyporządkowuje więc kolejne rodzaje dotychczas zbudowanych urządzeń napędowych do odpowiedniego czynnika roboczego (tj. do wiersza tablicy) oraz do odpowiedniej kategorii napędów (tj. do kolumny tablicy). Wzajemne uszeregowanie poszczególnych generacji napędów następuje według kompleksowości eksploatowanych atrybutów danego czynnika (tj. pierwsza generacja (1) eksploatuje tylko oddziaływania siłowe, druga (2) - oddziaływania siłowe i inercję, trzecia (3) zaś - oddziaływania siłowe, inercję i energię wewnętrzną). W obrębie każdej generacji wyróżniono dwie pary bliźniaczych urządzeń zwanych silnikiem i pędnikiem, eksploatujących te same cechy danego czynnika roboczego. Należy zwrócić uwagę, że w przypadkach gdy dane urządzenie budowane jest w wielu wersjach konstrukcyjnych, odmianach lub zastosowaniach, tylko pierwszą lub najbardziej reprezentacyjną jego wersję ujęto w tablicy (np. silnik parowy, turbina parowa, czy turbina gazowa wykorzystują te same atrybuty czynnika roboczego, stąd należą one do tego samego etapu rozwojowego). Jeszcze jedna "Tablica Cykliczności" pokazana jest jako tablica K1.

KOMORA OSCYLACYJNA

Motto: "Przyszłość zawsze wyrasta na przeszłości."

Wyobraźmy sobie małą kryształową kostkę stanowiącą nowe urządzenie do produkcji super-silnego pola magnetycznego. Wyglądałaby ona jak idealnych kształtów kryształ, lub jak ozdobny sześcian precyzyjnie wyszlifowany ze szkła i ukazujący swe fascynujące wnętrze poprzez przeźroczyste ścianki. Przy wielkości nie większej od poręcznej kostki Rubika wytwarzałaby ona pole setki tysięcy razy przewyższające pola produkowane dotychczas na Ziemi, włączając w to pola najsilniejszych współczesnych dźwigów magnetycznych czy najpotężniejszych elektromagnesów w laboratoriach naukowych. Gdybyśmy kostkę taką wzięli do ręki, wykazywałaby ona zdumiewające własności. Przykładowo mimo swych niewielkich rozmiarów byłaby ona niezwykle "ciężka" i przy jej przesterowaniu na pełny wydatek magnetyczny nawet najsilniejszy atleta nie byłby w stanie jej udźwignąć. Jej "ciężar" wynikałby z faktu, iż wytwarzane przez nią potężne pole magnetyczne powodowałoby jej przyciąganie w kierunku Ziemi i stąd do jej rzeczywistego ciężaru dodawałaby się wytworzona w ten sposób siła jej oddziaływań magnetycznych z polem ziemskim. Byłaby ona też oporna na nasze próby obracania i podobnie jak igła magnetyczna kompasu zawsze starałaby się zwrócić w tym samym kierunku. Gdybyśmy jednak zdołali ją obrócić w położenie dokładnie odwrotne do tego jakie sama skłonna byłaby przyjmować, wtedy ku naszemu zdumieniu zaczęłaby nas unosić w powietrze.

Podczas oglądania z niewielkiej odległości ta kryształowa kostka ukazałaby w swym wnętrzu niezliczone iskry elektryczne zamrożone w bezustannym migotaniu. Przy ich dokładniejszych oględzinach zauważylibyśmy, iż iskry te obiegają kostkę naokoło, "ślizgając" się wzdłuż powierzchni wewnętrznej jej czterech przeźroczystych ścianek bocznych. (Pozostałe dwie ścianki czołowe tego sześciennego kryształu stanowiłyby wyloty/bieguny magnetyczne wytwarzanego przez niego pola.) Przeskoki każdej z isker następowałyby tylko pomiędzy dwoma naprzemianległymi ściankami kostki. Ponieważ części drogi owych niezliczonych isker wzajemnie nachodziłyby na siebie, w sumie tworzyłyby one rodzaj "iskrowego wiru" jaki niezwykle szybko obiegałby wokół osi magnetycznej urządzenia. Wir ten nie zakreślałby jednak kolistych trajektorii jak to czynią wszystkie normalnie rotujące zjawiska, lecz jego iskry skakałyby wzdłuż obwodu kwadratu. Z kolei rotowanie tego obiegającego po kwadracie wiru iskrowego wytwarzałoby potężne pole magnetyczne w sposób podobny jak to czyni przepływ prądu wzdłuż obwodu selenoidu.

Powyższy opis ujawnił nam wygląd i działanie komory oscylacyjnej która stanowi przedmiot rozważań niniejszego rozdziału. Wynika z niego że nazwa "komora oscylacyjna" (po angielsku "Oscillatory Chamber") przyporządkowana została nieznanemu wcześniej na Ziemi urządzeniu realizującemu całkowicie nową zasadę wytwarzania pól magnetycznych, jaką miałem honor osobiście wynaleźć i rozpracować. Zasada ta wykorzystuje zjawisko rotowania czterosegmentowego łuku elektrycznego naokoło wewnętrznego obwodu czterech ścianek komory sześciennego. Łuk ten uformowany zostaje z dwóch pęków oscylujących isker elektrycznych skaczących w kierunkach wzajemnie prostopadłych do siebie, których cztery kolejne przeskoki formują cztery boki kwadratu. Przeskoki te następują pomiędzy naprzemianległymi ściankami komory sześciennego o specjalnej konstrukcji. Komora ta właśnie z uwagi na swój kształt i zasadę działania nazwana została „komorą oscylacyjną”.

Komora oscylacyjna realizująca powyższą zasadę działania uformowana będzie jako sześcienna kostka/pudło wykonana z przezroczystego materiału i pusta w środku (tj. wypełniona jedynie gazem dielektrycznym pod małym ciśnieniem). Jej sześć ścianek wykonane będzie z płytek materiału izolacyjnego (np. szkła) zespolonego na krawędziach. Na dwóch parach przeciwległych jej ścianek wewnętrznych zainstalowane zostaną pakiety elektrod przewodzących. Elektrody te (bez żadnych dodatkowych urządzeń uzupełniających) pełnić będą funkcję dwóch obwodów oscylacyjnych z iskrownikami. Każdy z obu tych obwodów oscylacyjnych powstanie w wyniku wzajemnego oddziaływania elektrycznego na siebie dwóch pakietów elektrod zamontowanych na przeciwległych ściankach wewnętrznych owej sześcienniej komory oscylacyjnej (tj. powierzchnia przeciwstawnych elektrod dostarczy obwodowi wymaganej pojemności elektrycznej, ich wzajemny odstęp służyć będzie jako przerwa iskrowa, zaś iskra elektryczna dostarczy wymaganej indukcyjności).

Realizacja przez komorę oscylacyjną jej unikalnej zasady wytwarzania pola magnetycznego, podsumowana w wielkim skrócie, jest następująca. Pakiety elektrod położonych na przeciwstawnych ściankach komory ładowane są przeciwnymi ładunkami elektrycznymi. Ładunki te próbują się wyrównać formując pęki isker oscylujących pomiędzy tymi przeciwstawnymi ściankami. Ponieważ dwa pęki takich oscylujących isker zmuszane są do naprzemiennie zsynchronizowanych przeskoków wzdłuż wewnętrznych powierzchni czterech ścianek komory oscylacyjnej, ich efektem końcowym będzie uformowanie łuku elektrycznego rotującego po obwodzie kwadratu. Łuk ten wytwarzał będzie wymagane pole magnetyczne. Powyższa realizacja działania komory oscylacyjnej umożliwia osiągnięcie podwójnych korzyści. Z jednej strony eliminuje ona bowiem prawie wszystkie wady wrodzone występujące w konstrukcji elektromagnesów, jakie ograniczają wydatek wytwarzanego przez nie pola magnetycznego. Z drugiej zaś strony dostarcza ona komorze oscylacyjnej dodatkowych zalet jakie są źródłem unikalnych (tj. nieznanymi wcześniej w żadnym innym urządzeniu budowanym dotychczas na Ziemi) atrybutów operacyjnych tego urządzenia.

Całkowite wyeliminowanie wad elektromagnesów uzyskane jest dzięki następującym atrybutom komory oscylacyjnej:

1. Pełna neutralizacja sił elektromagnetycznych działających na ścianki tej komory.
2. Pozostawienie użytkownikowi wyboru czasu zasilania oraz ilości energii dostarczanej do tej komory. Tj. każda porcja energii, niezależnie ile jej jest i kiedy zostaje ona dostarczona, jest przechwytywana przez tą komorę, przechowywana, zamieniana na pole magnetyczne, oraz uwalniana kiedy staje się potrzebna.
3. Odzyskiwanie i zamiana z powrotem na elektryczność całości energii rozpraszanej przez iskry.
4. Kierowanie niszczycielskich następstw zakumulowania ogromnych ilości energii w sposób jaki wzmacnia zasadę działania komory i nie niszczy jej materiału.
5. Niezależność mocy urządzeń sterujących od mocy wytwarzanego pola magnetycznego (tj. słaby sygnał sterujący powoduje zmiany w ogromnie silnym polu magnetycznym wytwarzanym przez tą komorę).

Komora oscylacyjna wykazuje także następujące zalety jakie nie znane były dotychczas w żadnym urządzeniu budowanym przez człowieka:

- A. Zdolność do zaabsorbowania i przechowania nieograniczonej ilości energii.
- B. Pełna kontrola nad wszystkimi atrybutami i parametrami wytwarzanego pola, uzyskiwana bez zmiany w całkowitej ilości energii zakumulowanej w tej komorze.
- C. Wytwarzanie rodzaju pola magnetycznego jakie nie przyciąga ani nie odpycha obiektów ferromagnetycznych (tj. jakie zachowuje się jak owo hipotetyczne "pole antygrawitacyjne", nie zaś jak pole magnetyczne).

D. Wielowymiarowa transformacja energii (np. elektryczność - pole magnetyczne - ciepło) która umożliwia komorze oscylacyjnej przejęcie funkcji niemal wszystkich innych konwencjonalnych urządzeń konwersji energii (np. elektromagnesów, transformatorów, generatorów, akumulatorów, baterii, silników spalinowych, grzejników, klimatyzatorów powietrza, oraz wielu więcej).

Końcowym wynikiem takiego opracowania konstrukcji i działania komory oscylacyjnej jest, iż po zbudowaniu urządzenie to będzie zdolne do zwiększania swego wydatku magnetycznego do teoretycznie Nielimitowanego poziomu. Praktycznie oznacza to, że owo źródło potężnego pola magnetycznego będzie pierwszym urządzeniem zbudowanym na Ziemi, którego wydatek magnetyczny umożliwi mu przekroczenie strumienia startu, a co za tym idzie samoczynne wzniesienie się w powietrze jedynie w efekcie odpychającego oddziaływania wytwarzanego przez siebie pola z polem magnetycznym Ziemi, Słońca lub Galaktyki. Komora oscylacyjna będzie więc naszym pierwszym "magnesem zdolnym do wzlotu".

Z uwagi że komora oscylacyjna jest urządzeniem napędowym (pędnikiem) każdego statku międzygwiazdowego (jako przykład przeglądaj też opisy prognostyczne z podrozdziału M6), że jest ona najistotniejszym elementem składowym prawie wszystkich wyrobów budowanych przez co bardziej zaawansowane cywilizacje (patrz przegląd jej zastosowań zaprezentowany w podrozdziale C9), oraz że jej zbudowanie zadecyduje o awansowaniu naszej cywilizacji ze szczebla planetarnego do szczebla międzygwiazdowego (co dokładniej objaśniono w podrozdziale M6), w niniejszym rozdziale zdecydowałem się przytoczyć w miarę szczegółowe jej omówienie. Niemniej uwadze czytelników zainteresowanych w historii rozwoju tej nowej myśli technicznej, polecane są również starsze monografie [1/3], [1/2], [3/2], [3] i [2], które także zawierały rozdziały jak niniejszy, poświęcone wcześniejszym prezentacjom komory oscylacyjnej oraz jej najważniejszych zastosowań.

C1. Dlaczego niezbędnym jest zastąpienie elektromagnesów przez komorę oscylacyjną

Obserwując osiągnięcia naszej wiedzy i techniki w jednej dziedzinie, np. przemyśle spożywczym, bez zastanowienia zakładamy, że nasz postęp jest równie efektowny we wszystkich kierunkach. Tymczasem istnieją działy techniki, gdzie nie nastąpił prawie żaden postęp od prawie dwóch stuleci i gdzie ciągle drepjemy w kółko w tym samym miejscu. Aby uświadomić sobie jeden z najbardziej powszechnie spotykanych przykładów takiego zastoju, zadajmy teraz pytanie: "Jakiż to postęp dokonany został ostatnio w zakresie zasad wytwarzania sterowalnych pól magnetycznych?" Zaskakująco, odpowiedź jest: "żaden". W dobie eksploracji Marsa do wytwarzania pola magnetycznego ciągle wykorzystujemy dokładnie tą samą zasadę, jaka wykorzystywana była w tym celu przed ponad 170 lat, tj. zasadę odkrytą w 1820 roku przez duńskiego profesora Hans'a Oersted'a i polegającą na wykorzystaniu efektów magnetycznych prądu elektrycznego przepływającego przez zwoje przewodnika. Urządzenie wykorzystujące tą zasadę, nazywane "elektromagnesem", jest obecnie jednym z najbardziej archaicznych wynalazków ciągle w powszechnym użyciu z powodu braku lepszego rozwiązania. Aby zrozumieć jak przestarzałe jest działanie elektromagnesu wystarczy posłużyć się następującym przykładem: gdyby nasz postęp w rozwoju urządzeń napędowych równał się postępowi w rozwoju urządzeń do wytwarzania pól magnetycznych, wtedy naszym jedynym wehikułem ciągle pozostawałaby lokomotywa parowa.

Elektromagnesy posiadają cały szereg wad wrodzonych. Wady te uniemożliwiają podniesienie ich wydatku ponad określony, i to stosunkowo niski, poziom. Ich usunięcie nie jest możliwe w żaden sposób, ponieważ wynikają one z samej zasady działania tych urządzeń. Poniżej dokonany zostanie przegląd najważniejszych z owych **nieusuwalnych wad** elektromagnesów. Ich bardziej dokładne omówienie przytoczone jednak będzie w podrozdziale C6 poświęconym prezentacji zasad, na jakich każda z nich wyeliminowana została w działaniu komory oscylacyjnej.

#1. Elektromagnesy formują elektromagnetyczne **siły odchylające**, jakie napinają ich zwoje w kierunku promieniowym, starając się rozerwać te zwoje na strzępy. Siły te formowane zostają w rezultacie wzajemnego oddziaływania pomiędzy polem magnetycznym

wytwarzanym przez dany elektromagnes, a zwojami przewodnika jaki wytworzył to pole. Pole to, zgodnie z działaniem "reguły lewej ręki" często zwanej także "efektem silnika", stara się wypchnąć zwoje wytwarzającego je przewodnika ze swego zasięgu. Elektromagnetyczne siły odchylające uformowane w ten sposób są więc identycznego rodzaju jak te wykorzystywane w zasadzie działania silników elektrycznych. Aby zabezpieczyć elektromagnes przed rozerwaniem na strzępy, owym wewnętrznym elektromagnetycznym siłom odchylającym musi przeciwstawić się jakaś fizyczna konstrukcja zewnętrzna. Konstrukcja ta balansuje swoją mechaniczną wytrzymałością siły odchylające wynikające z wydatku danego elektromagnesu. Konstrukcja owa oczywiście niezwykle zwiększa wagę każdego silniejszego elektromagnesu. Więcej jednak, jeśli przepływ prądu w elektromagnesie przewyższy określony poziom, wtedy owe siły odchylające wzrastają do takiej wartości, iż żadna fizyczna konstrukcja nie jest już w stanie im się oprzeć, dlatego też powodują one eksplodowanie zwojów danego elektromagnesu. W ten sposób zbyt duże zwiększenie wydatku dowolnego elektromagnesu zwykle kończy się jego **samo-zniszczeniem poprzez eksplodowanie**. Takie eksplozje elektromagnesów są dosyć częstym zdarzeniem w laboratoriach badawczych, stąd co potężniejsze z owych urządzeń montowane są w specjalnych bunkrach wyciszających skutki ich ewentualnej eksplozji.

#2. Wymogiem elektromagnesów jest **bezustanne zaopatrywanie w energię** elektryczną jeśli produkowane przez nie pola muszą posiadać kontrolowalne parametry (tj. jeśli parametry ich pola są zmienialne zgodnie z wymaganiami użytkowników). Jeśli takie bezustanne zaopatrywanie w energię elektryczną zostanie nagle odcięte, sterowalność ich pola magnetycznego także ulega zakończeniu. Powyższe wymaganie nałożone na sterowalność pola elektromagnesów powoduje, że podczas produkcji potężnych pól magnetycznych, pojedynczy elektromagnes konsumuje wydatek całej elektrowni.

#3. Elektromagnesy powodują liczące się **straty energii**. Prąd elektryczny przepływający przez zwoje konwencjonalnego elektromagnesu wyzwala ogromne ilości ciepła (patrz Prawo Joule'a dotyczące nagrzewania prądem elektrycznym). To ciepło nie tylko że pomniejsza efektywność energetyczną produkcji pola, ale także - kiedy energie pola są wysokie, powoduje ono topienie się zwojów elektromagnesu.

Użycie materiałów nadprzewodzących do wykonania zwojów elektromagnesu eliminuje wprawdzie nagrzewanie się jego materiału w efekcie przepływu prądu. Jednakże równocześnie wprowadza ono innego rodzaju straty energii wynikające z konieczności utrzymywania bardzo niskiej temperatury zwojów elektromagnesu nadprzewodzącego. Oczywiście takie utrzymywanie temperatury wiąże się z bezustanną konsumpcją energii jaka zmniejsza efektywność wynikową danego elektromagnesu. Powinno także tu zostać podkreślone, że pole magnetyczne o wysokiej gęstości eliminuje efekt nadprzewodnictwa i stąd przywraca oporność elektryczną do zwojów. Dlatego też elektromagnesy nadprzewodzące są tylko w stanie wytwarzać pola leżące poniżej owej wartości progowej powodującej nawrót ich oporności elektrycznej.

#4. Elektromagnesy są podatne na **zużycie elektryczne**. Konfiguracja geometryczna elektromagnesów jest tak uformowana, że kierunek największych sił pola elektrycznego nie pokrywa się z ułożeniem przewodnika w zwoje (tj. siły tego pola starają się powodować przepływ prądu w poprzek uzwojeń, podczas gdy ułożenie warstewek izolacyjnych wymusza ten przepływ wzdłuż zwoi po spirali). To z kolei skierowuje niszczące działanie energii elektrycznej na izolacje zwojów elektromagnesu. Po upływie określonego czasu energia ta powoduje więc przebicie elektryczne izolacji jakie inicjuje zniszczenie całego tego urządzenia (tj. wywołuje spięcie elektryczne w uzwojeniach elektromagnesu które następnie topi zwoje i niszczy cały elektromagnes).

#5. Elektromagnesy uniemożliwiają **sterowanie** swoim działaniem za pomocą słabych sygnałów sterujących. Parametry wytwarzanego przez nie pola magnetycznego mogą być zmieniane tylko poprzez zmiany w mocy zasilającego prądu elektrycznego. Dlatego też kontrolowanie elektromagnesu wymaga użycia tych samych mocy jak moce niezbędne do wytwarzania pola magnetycznego.

Jedyną drogą do wyeliminowania pięciu powyższych najistotniejszych wad wrodzonych elektromagnesów jest zastosowanie do wytwarzania pola całkowicie odmiennej zasady działania. Zasada taka, jaką miałem honor osobiście wynaleźć, zostanie zaprezentowana w dalszych częściach tego rozdziału. Ponieważ zasada ta wykorzystuje mechanizm oscylacyjnych wyładowań elektrycznych następujących we wnętrzu komory sześcienniej, nazwana ona została "komorą oscylacyjną".

Zasada działania komory oscylacyjnej nie posiada powyżej opisanych wad i ograniczeń elektromagnesów uniemożliwiających zwiększanie ich wydatku (sposób eliminacji w niej owych wad opisano w podrozdziale C6). Ponadto, zapewnia ona bardziej proste oraz efektywne wytwarzanie i użytkowanie, długi okres użytkowania bez konieczności dokonywania napraw, niezwykle wysoką wartość stosunku wydatek-do-ciężaru, oraz szeroki zakres zastosowań komory (np. jako urządzenia napędowego, akumulatora energii, źródła pola magnetycznego, itp. - patrz tablica C1). Wyjaśnienia jakie nastąpią (szczególnie te z podrozdziału C7) opiszą mechanizm umożliwiający osiągnięcie wszystkich tych dodatkowych zalet. Niewystępowanie więc w komorze oscylacyjnej wszystkich wad wrodzonych elektromagnesów, w połączeniu z licznymi zaletami operacyjnymi, uzasadniają celowość szybkiego podjęcia jej budowy, tak aby to niezwykle urządzenie już wkrótce mogło zastąpić obecnie używane elektromagnesy.

C2. Historia komory oscylacyjnej

Jak każde urządzenie techniczne opracowywane przez dłuższy okres czasu, również komora oscylacyjna posiada już swoją historię. Co może najbardziej fascynować, to że obecnie wszyscy współuczestniczymy w tworzeniu tej historii. Komora oscylacyjna jest prawdopodobnie jednym z niewielu urządzeń, którego wynalezienie dokonane zostało "na żądanie". Stąd dla badaczy z niektórych dyscyplin analiza jej historii może prowadzić do interesujących ustaleń. Ponadto opisanie tu sposobu zesyntezowania tego wynalazku uświadomi też czytelnikowi, że za suchymi opisami technicznymi i bezosobową matematyką niniejszego rozdziału kryje się pasjonująca historia ludzkich zmagania i intelektualnego wyzwania rzuconego naturze. Przytoczmy więc tu krótki zarys historii wynalezienia tego urządzenia. Historia komory oscylacyjnej doskonale uzupełnia i poszerza generalną historię niniejszej monografii, zaprezentowaną w podrozdziale A4. Na wypadek jednak gdyby czytelnik nie czytał jeszcze podrozdziału A4, poniżej powtórzone opisy wszystkich tych wydarzeń, jakie wywarły kluczowy wpływ na rozwój i ewolucję komory oscylacyjnej. Jako że również i w historii komory oscylacyjnej wyraźnie daje się wyodrębnić kilka istotnych "kamieni milowych", historia ta zestawiona tutaj będzie w ujęciu etapowym, po jej uprzednim posegmentowaniu na takie istotne "kamienie milowe (#)".

#1. Stworzenie potrzeby wynalezienia komory oscylacyjnej. Jak to już wyjaśnione zostało w podrozdziałach A1 i A4, wszystko zaczęło się 1972 roku. Prowadziłem wtedy serię wykładów dla studentów Politechniki Wrocławskiej, poświęconych "wybranim zagadnieniom systemów napędowych". W czasie przygotowywania jednego z tych wykładów odkryłem, że w działaniu urządzeń napędowych zbudowanych dotychczas na Ziemi istnieje zdumiewająca symetryczność. Symetryczność tą nazwałem później "Prawem Cykliczności". Najlepiej ją wyrazić za pomocą tzw. "Tablicy Cykliczności". Pierwszy opis owej Tablicy Cykliczności opublikowany został w artykule **[1C2]** "Teoria rozwoju napędów", jaki ukazał się w polskim czasopiśmie Astronautyka, numer 5/1976, str. 16-21. Przykład jej obecnej postaci pokazany został jako tablica B1. Natomiast jej sporządzanie streszczone zostało w podrozdziałach B1 i K1 tej monografii, oraz dodatkowo wyjaśnione w innych publikacjach jak [1], [1a], [3], [3/2], [6] i [6/2] wyszczególnionych we wykazie literatury uzupełniającej treść niniejszej monografii (patrz rozdział Y).

Tablice Cykliczności stanowią odmianę "Tablicy Mendelejewa", tyle tylko że zamiast dla pierwiastków chemicznych opracowywane są one dla urządzeń technicznych.

Uwidaczniają one, że kolejne odkrycia tego samego rodzaju urządzeń (np. silników i pędników) układają się w symetryczny i powtarzalny wzór podobny do wzoru wykazywanego przez kolejne pierwiastki chemiczne Tablicy Mendelejewa. Tablice Cykliczności dają się sporządzić dla prawie wszystkich rodzajów urządzeń technicznych, nie zaś tylko dla napędów. Z kolei analizując takie uprzednio sporządzone już tablice, możliwym się staje prognozowanie przyszłego rozwoju urządzeń opisanego nimi rodzaju. Prognozowanie to zezwala nie tylko na przewidywanie jakie dalsze urządzenia danego rodzaju nadal czekają na swego wynalazcę i budowniczego, ale także jaka będzie przyszła zasada działania owych dotychczas jeszcze nie wynalezionych urządzeń. Poprzez studiowanie wzajemnych odstępów czasowych pomiędzy datami zbudowania już istniejących urządzeń, Tablice Cykliczności pozwalają też wnioskować o dynamice działalności wynalazczej dotyczącej urządzeń danego przeznaczenia. To z kolei umożliwia wyznaczanie przybliżonych dat budowy przyszłych generacji tych urządzeń.

Analizując pierwszą opracowaną przez siebie Tablicę Cykliczności (pokazaną tu jako tablica B1), odkryłem z niej, że wkrótce na naszej planecie powinna zostać zbudowana nowa generacja wehikułów latających, jakie później nazwałem magnokrafty. Zasada działania tych wehikułów stanowić będzie rozszerzenie działania asynchronicznych silników elektrycznych. Jako swoje pędniki magnokrafty wykorzystywały będą potężne "magnesy" (tj. pędniki magnetyczne), które formowały będą siły napędowe na wskutek przyciągania i odpychania wytwarzanych przez siebie pól magnetycznych z polami magnetycznymi otoczenia (tj. polami Ziemi, Słońca, lub Galaktyki). Praktycznie więc, pod względem użytej zasady działania magnokrafty stanowiły będą następcę współczesnych silników elektrycznych. Tyle tylko, że dla wytwarzania ruchu zamiast pola ze stojana, wykorzystywały one będą pole magnetyczne swego otoczenia. W 1980 roku opublikowałem pierwsze szczegóły techniczne magnokraftów w artykule **[2C2]** "Budowa i działanie statków kosmicznych z napędem magnetycznym" jaki ukazał się w Przeglądzie Technicznym Innowacje (nr 16/1980, str. 21-23). Narodziny idei tego statku wprowadziły z kolei potrzebę aby opracowane dla niego zostało jakieś urządzenie napędowe. Urządzeniem tym potem okazała się być komora oscylacyjna.

#2. Uświadomienie sobie, że komora oscylacyjna musi być wynaleziona osobiście przeze mnie. Wynalezienie, rozpracowanie i późniejsze upowszechnienie budowy i działania magnokraftu, wprowadziło do użycia pojęcie "pędnika magnetycznego". Zaczniemy więc od zdefiniowania tego pojęcia, bowiem rzutuje ono na mechanizm stojący za uświadomieniem mi konieczności osobistego wynalezienia komory oscylacyjnej.

"Pędnikiem magnetycznym" nazywane jest źródło silnego pola magnetycznego (tj. rodzaj 'magnesu'), którego wydatek przekracza wartość progową zwaną 'strumieniem startu'. Oczywiście powyższa definicja jest jedynie zrozumiała, jeśli wiemy czym właściwie jest ów strumień startu. Przytoczmy więc tutaj wypracowaną w podrozdziale F5.1 definicję tego pojęcia. **"Strumieniem startu (F_s)"** nazwana jest taka wartość progowa pola magnetycznego wytwarzanego przez niezwykle potężne jego źródło, jaka po odpychającym zorientowaniu tego źródła względem naturalnego pola magnetycznego planety Ziemia, byłaby w stanie wytworzyć siły magnetycznego odpychania zdolne wznieść to źródło w przestrzeń kosmiczną. Strumień startu jest więc rodzajem magnetycznego odpowiednika dla pierwszej prędkości kosmicznej. Każdy bowiem 'magnes' jaki wytworzy ten strumień, będzie w stanie sam się wznieść w powietrze, jeśli tylko ktoś zorientuje go odpychająco względem ziemskiego pola magnetycznego. (Takie odpychające zorientowanie względem pola otoczenia polega na ustawieniu tego 'magnesu' w pozycji dokładnie odwrotnej do pozycji, jaką on sam byłby skłonny przyjmować, gdyby udzielić mu swobody obrotu podobnej do tej posiadanej przez igły kompasów magnetycznych.) Oczywiście, Teoria Magnokraftu umożliwia precyzyjne wyliczenie wartości tego strumienia. Ja już dokonałem odpowiednich obliczeń (patrz podrozdział F5.4) i ustaliłem, iż wyznaczona dla obszaru Polski wartość "strumienia startu" wynosi $F_s = 3.45$ [Wb/kg]. Aby więc zbudować pierwszy pędnik magnetyczny, opracowany musi zostać jakiś rodzaj sterowalnego magnesu który byłby zdolny do wytworzenia tak dużego wydatku. Niestety, jak to już wyjaśniłem w podrozdziale C1, przy użyciu zasady działania elektromagnesów, wytworzenie wydatku o takim poziomie nie jest możliwe.

Rozpoczęcie prac rozwojowych nad magnokraftem wprowadziło więc dwie niezwykle istotne konsekwencje. Pierwszą i naistotniejszą z nich było, że uświadomiło mi to iż rozwijanie budowy i działania magnokraftu wymaga wynalezienia zupełnie nowej zasady wytwarzania pola magnetycznego. Wszakże żadna z dotychczas istniejących takich zasad nie pozwalała na przekroczenie wartości strumienia startu. Owa nowa zasada jaka wówczas ciągle oczekiwała na wypracowanie, musiała więc być nastawiona na spełnienie podstawowego wymagania nakładanego na pędnik magnetyczny. Mianowicie, jej wydatek magnetyczny musiał przekraczać krytyczną wartość "strumienia startu". W ten sposób magnokraft uświadomił mi konieczność osobistego rozpoczęcia prac wynalazczych nad rewolucyjnym urządzeniem technicznym, które później nazwane zostało "komorą oscylacyjną".

Drugą konsekwencją rozpracowania idei magnokraftu było, że rozpracowanie to ujawniło **podstawowe wymaganie** (warunek operacyjny) nałożone na źródło pola magnetycznego, aby źródło to mogło zostać użyte jako pędnik magnetyczny. Warunkiem tym jest, że wydatek takiego źródła musi przekroczyć krytyczną wartość "strumienia startu". Ja byłem całkowicie świadom tego warunku od pierwszej chwili swoich prac rozwojowych nad magnokraftem. Wszakże matematyczne zaprezentowanie tego warunku nastąpiło już w pierwszej publikacji [1C2] poświęconej teorii napędów magnetycznych. Niestety, w początkowym stadium swych badań nie wiedziałem jakie urządzenie byłoby w stanie sprostać temu wymaganiu. Wiadomo bowiem, iż elektromagnesy używane obecnie do produkcji najpotężniejszych z pól magnetycznych dostępnych naszej cywilizacji, posiadają wady konstrukcyjne jakie uniemożliwiają nawet zbliżenie się ich wydatku do owego krytycznego strumienia startu. (Owe wady elektromagnesów omawiane są w podrozdziałach C1 i C6.)

#3. Znalezienie kierunku poszukiwań twórczych. Jak to podkreśla niniejszy rozdział, od pierwszej chwili skryształizowania się idei magnokraftu byłem całkowicie świadom, że dotychczasowe urządzenia wytwarzające pole magnetyczne nie będą w stanie dostarczyć wydatku przekraczającego strumień startu. Fakt naszej nieznajomości takiego urządzenia był też powodem zaciekłych ataków magnokraftu ze strony wielu jego przeciwników. Stąd dla mnie znalezienie konceptu takiego urządzenia stanowiło problem oczekujący najpilniejszego rozwiązania. Przemyślałem więc nad nim niemalże bez ustanku. Niestety postawiony cel zdawał się niezwykle trudny do osiągnięcia. Wszakże przez prawie dwa stulecia całe generacje naukowców i wynalazców bezskutecznie głowiły się nad znalezieniem jakiegoś lepszego od elektromagnesów sposobu wytwarzania pola magnetycznego. Niedługo przed opuszczeniem Polski, podczas jednego z wypoczynkowych pobytów w Karpaczu zimą 1981 roku, zaobserwowałem obciążoną ciężarówkę jaka z widoczną trudnością wdrapywała się na strome zbocze góry. Obserwacja tej ciężarówki uświadomiła mi, że działanie poszukiwanego przeze mnie urządzenia, musi być oparte na jakiejś formie zamiany ruchu oscylacyjnego w ruch jednostajny. (Podobnie jak w silniku ciężarówki ruch oscylacyjny tłoka zamieniany jest w jednostajny ruch obrotowy jej kół.) Nie może więc ono być oparte na ciągłym przepływie energii, jak to jest w przypadku elektromagnesów. To przełomowe ustalenie nadało z kolei prawidłowy kierunek moim dalszym poszukiwaniom zasady działania komory oscylacyjnej.

#4. Ukierunkowana synteza wynalazku komory. Po uświadomieniu sobie, że komora oscylacyjna musi dopiero zostać wynaleziona, oraz znalezieniu kierunku w jakim kryje się jej rozwiązanie, rozpocząłem systematyczne poszukiwania twórcze mastawione na dokonanie "ukierunkowanej syntezy nowego wynalazku". Istotnym elementem takiej ukierunkowanej (po angielsku "goal oriented") syntezy jest, że produkt końcowy jaki nasz umysł poszukuje, jest ściśle zdefiniowany i obłożony licznymi warunkami operacyjnymi. Jest to więc wyższy szczebel działalności wynalazczej, ponieważ w normalnych przypadkach wynalazki polegają na wpadaniu na idee zupełnie niezwiązane z kierunkiem w jakim wynalazca podąża, czy z rozwiązaniem jakiego on właśnie poszukuje. (Typowe wynalazki następują więc zgodnie z popularnym powiedzeniem "mamy już lekarstwo, poszukajmy teraz odpowiedniej choroby".)

Pomimo świadomości, że poszukiwane urządzenie musi bazować na ruchu oscylacyjnym zamienianym na ciągly przepływ energii, mój umysł nadal był jednak więziony stereotypami jakie wówczas panowały, a jakie sugerowały że urządzenie wytwarzające pole magnetyczne musi mieć formę pierścienia lub okrągłej cewki. Z uwagi na ten stereotyp, idei dla swego urządzenia poszukiwałem wśród różnorodnych już istniejących konstrukcji zabezpieczających pierścieniowy obieg ładunków elektrycznych, takich jak przykładowo TOKAMAK. Pracując nad różnymi możliwymi conceptami, analizowałem ogromną liczbę różnorodnych urządzeń technicznych, jakich działanie łączy w sobie drgania elektryczne, plazmę lub iskry, ruch naładowanych cząsteczek, itp. W ten sposób w swoim umyśle stopniowo zgromadziłem wszystkie elementy układanki obecnie zwanej "komora oscylacyjna". Tyle tylko, że początkowo elementy te ciągle podobne były do wymieszanych z sobą kawałków obrazkowej łamigłówki. Koniecznym więc był proces jakiegoś ich dopasowania w jedną całość.

#5. Znalezienie klucza dla zasady działania komory. Dopasowanie do siebie elementów łamigłówki zawartej w mojej głowie nastąpiło w nocy z 2-go na 3-ci stycznia 1984 roku. Korzystając wówczas z letnich wakacji w Nowej Zelandii, wybrałem się z kilku powodów do Christchurch. Pomimo zajmowania się innymi sprawami, mój umysł cały czas pracował jednak nad znalezieniem rozwiązania dla nękającego mnie problemu. Nieco po północy, gdy myślałem nad swoim problemem leżąc w łóżku w stanie prawie że półsnu, rozwiązanie nagle zostało zeszyntezowane w moim umyśle. **Kluczem** okazał się fakt, iż poszukiwane urządzenie musi przyjąć formę kostki sześciennnej, nie zaś pierścienia. Ciągle doskonale pamiętam, że fakt końcowego znalezienia tego długo poszukiwanego rozwiązania okazał się tak podniecający, że mimo wstania z łóżka w celu dokonania natychmiastowych notatek, rysunków i sprawdzeń, nie byłem w stanie utrzymać długopisu w dygoczącej z emocji ręce.

W ten więc sposób w pierwszych dniach 1984 roku, po kilku latach nieustannych poszukiwań, zeszyntezowałem w swoim umyśle concept nowego urządzenia, które będzie w stanie wytworzyć wydatek magnetyczny przekraczający strumień startu, nie eksplodując przy tym ani nie rozpadając się w kawałki. Owo niezwykle urządzenie nazwałem "komora oscylacyjna". Jej wynalezienie było jednak jedynie pierwszym krokiem na długiej i trudnej drodze wiodącej do jej realizacji.

#6. Pierwsze publikacje oraz fala entuzjastycznych eksperymentów nad jej realizacją. Zaraz po zeszyntezowaniu komory oscylacyjnej, upowszechniłem jej opisy w całym szeregu publikacji napisanych w trzech językach (angielskim, polskim, oraz niemieckim). Publikacje te dostępne były w czterech różnych krajach, tj. Nowej Zelandii, Polsce, USA i Niemczech Zachodnich - patrz podrozdział C10. Pierwsza z nich ukazała się jeszcze w styczniu 1984 roku - patrz [1C](a). Łatwość dostępu do opisów tego urządzenia w połączeniu z jego niezwykle atrakcyjnością spowodowała spore w nim zainteresowanie. Cały szereg indywidualnych amatorów i małych przedsiębiorstw rozpoczęło prace rozwojowe nad zbudowaniem pierwszego prototypu komory oscylacyjnej. Oczywiście, jak to ostatnio coraz częściej bywa z nowymi ideami w strategicznych dla ludzkości dziedzinach (patrz podrozdział VB5.1.2), lista zainteresowanych stron nie zawierała nawet jednego reprezentanta instytucji jakie powinny czuć się odpowiedzialne za postępy w urządzeniach do wytwarzania pola magnetycznego. Mianowicie, rozwoju komory oscylacyjnej nie podjęło żadne laboratorium naukowe pracujące nad polami magnetycznymi - na przekór zachęty i dokładnych opisów dostarczanych przeze mnie do użytku dużej liczby takich instytucji. Większość amatorów zainteresowanych w budowie komory oscylacyjnej pochodziło z Niemiec Zachodnich, Szwajcarii, Austrii i Polski.

#7. Elektrody igłowe, jako przykład rozwiązania technicznego dla pierwszego poważnego problemu konstrukcyjnego. Jak to daje się przewidzieć z opisu komory oscylacyjnej, zbudowanie prototypu tego urządzenia stanowi trudne zadanie. Dlatego też, jeden po drugim, większość początkowych budowniczych pomału dała za wygraną i wycofała się z projektu. Wśród tych których nie zraziły istniejące trudności i kontynuowali badania, był także nasz rodak, który niestety po pierwszych doświadczeniach ze "staniem się sławnym"

prosił mnie, aby ponownie nie publikować już jego nazwiska. W maju 1987 roku przesłał on mi zdjęcie swego modelu komory oscylacyjnej, jakie uchwyciło pęk iskier elektrycznych w ruchu rotującym. Fotografia tego modelu pokazana została na rysunku C13.

Problem techniczny, który już na samym początku zniechęcił większość początkowych budowniczych komory, zilustrowany został na rysunku C2 (a). Podążając bowiem za dostępnymi dla nich opisami komory, w pierwszych jej modelach budowanych przez siebie, starali się oni zastosować płyto-kształtne elektrody, tak jak to pokazano na rysunku C1 (b). Jednakże, gdy takie elektrody zostają użyte, iskry zamiast przeskakiwać grzecznie jak to się od nich spodziewa wzdłuż drogi na rysunku C2 (a) oznaczonej jako S', raczej wolą podążać wzdłuż linii najmniejszego oporu i przeskakiwać wzdłuż drogi oznaczonej tam jako S". Różni badacze starali się rozwiązać ten problem na kilka odmiennych sposobów, zaczynając od umieszczania tych elektrod wewnątrz cel izolacyjnych w kształcie plastra pszczelego, a kończąc na pokrywaniu elektrod warstwą izolacyjną. Dopiero nasz rodak znalazł właściwe rozwiązanie. Poprzez podążanie za moimi wskazówkami w niniejszej monografii zaprezentowanymi w rozdziale S5, rozpoczął on studiowanie opisów Arki Przymierza. Wnioskiem końcowym do jakiego doszedł po tych studiach było, iż Arka nie zawierała w swym wnętrzu żadnych płyto-kształtnych elektrod. Jedynie czubki złotych gwoździ były wbite poprzez jej drewniane ścianki i wystawały po wewnętrznej stronie tych ścianek. Stąd zdecydował się on rozpocząć eksperymenty z igłowymi elektrodami. I to rozwiązanie okazało się działać w praktyce. Takie igły odpychają iskry przeskakujące w ich pobliżu, stąd iskry te nie są w stanie skrócić swojej drogi poprzez wnikanie do materiału elektrod. W ten sposób, prototyp komory oscylacyjnej jaki zamiast płytek wykorzystywał igłowe elektrody - jak to pokazano na rysunku C2 (b), był pierwszą realizacją zasady komory jaka eksperymentalnie wytworzyła uporządkowane pęki iskier. Prototyp ten dostarczył więc eksperymentalnego potwierdzenia, iż zasada działania komory oscylacyjnej jest możliwa do technicznego zrealizowania w formie działającego urządzenia, konkludując w ten sposób etap zerowy rozwoju komory (patrz podrozdział C8.2).

#8. Procedura badań rozwojowych nad komorą oscylacyjną. Dokonywana przeze mnie analiza przyczyn niepowodzeń i postępującego zniechęcania się różnych hobbystów i badaczy pracujących nad zbudowaniem prototypu komory oscylacyjnej ujawniła, iż jedną z najważniejszych z tych przyczyn jest brak wyraźnych wskazówek jak właściwie zabrać się do systematycznej realizacji tak kompleksowego projektu. Zdecydowałem się więc opracować rodzaj niezawodnej procedury badawczej, która nieuchronnie wiodłaby do sukcesu, gdyby ktoś znalazł w sobie wystarczająco dużo sił i motywacji aby ją skompletować. Nazwałem ją "procedurą małych kroczków". Dla ułatwienia jej zastartowania opracowałem też kilka inicjujących tematów badawczych (w niniejszej monografii zaprezentowanych w podrozdziale C8.3). Procedura ta, wraz z towarzyszącymi jej tematami inicjującymi, po raz pierwszy została opublikowana 27 stycznia 1994 roku w monografii [2], potem zaś 13 września 1994 roku powtórzona w monografii [2a]. W niniejszym opracowaniu jest ona przytoczona w podrozdziale C8.2.

#9. Komora oscylacyjna drugiej generacji. Wynalezienie i rozpracowanie w 1989 roku zasady działania baterii telekinetycznych naprowadziło mnie na koncept komór oscylacyjnych drugiej generacji. W niniejszej monografii komory te omówione zostały w podrozdziałach C4.1, C7.1.1, C7.1.2 i C7.2.2. Ich cechą charakterystyczną jest, że będą one zdolne do wytwarzania efektu telekinetycznego opisanego w podrozdziale H6.1 (szczególnie patrz podrozdziały H6.1 i L2). Efekt ten z kolei z jednej strony umożliwi wyposażonym w nie magnokraftom działanie w telekinetycznej konwencji lotu. (Czym jest owa telekinetyczna konwencja lotu opisano w podrozdziałach B1 i L1.) Z drugiej zaś strony, zezwoli on też komorom oscylacyjnym drugiej generacji na samowzbudzenie swoich oscylacji. W ten sposób będą one zdolne do pracy jako baterie telekinetyczne, samozapelniając się w ściśle kontrolowany sposób wymaganą ilością energii. Jako takie, działać one będą nie tylko jako pędniki telekinetyczne oraz akumulatory energii o ogromnej pojemności, ale także jako siłownie telekinetyczne produkujące i samoladujące się całą zużywaną przez siebie energię.

Generalny kierunek badań nad zrealizowaniem samoladujących się komór drugiej generacji wskazuje podrozdział K2.4. Niemniej czytelnicy upominani są aby przypadkiem samemu nie rozpoczynać eksperymentów nad komorami drugiej generacji, zanim wypracowane zostaną efektywne systemy zabezpieczające komory oscylacyjne przed przeładowaniem energią i przed eksplozowaniem. Wszakże ich ewentualne eksplozowanie mogłoby okazać się ogromnie niszczycielskie - patrz opisy z podrozdziału O5.2 oraz prezentacje w odrębnych monografiach z serii [5]. Eksplozowanie takie mogłoby wysadzić w powietrze nie tylko dom niefortunnego eksperymentatora, ale także i całe miasto w którym by się on znajdował.

Wypracowanie idei komory oscylacyjnej drugiej generacji wskazało też generalną zasadę podążanie po której umożliwiałoby będzie budowanie komór oscylacyjnych coraz wyższych generacji. Ekstrapolując tą zasadę, później rozpracowałem również podstawy funkcjonalne **komór oscylacyjnych trzeciej generacji** - patrz podrozdział C4.1. Pierwsze opublikowanie informacji na temat komór oscylacyjnych drugiej generacji nastąpiło 27 stycznia 1994 roku w podrozdziale C3.1 monografii [2].

#10. Odkrycie fal telepatycznych. W piątek 11 listopada 1994 roku podczas przerwy na lunch odkryłem czym właściwie są fale telepatyczne. Odkrycie to dokładnie opisałem już w podrozdziale A4. Zgodnie z nim "fale telepatyczne są dźwiękopodobnymi wibracjami przeciw-materii". (Zauważ, że zgodnie z treścią podrozdziału H5.2, wszelki ruch przeciw-materii manifestuje się w naszym świecie jako pole magnetyczne. Stąd fale telepatyczne można też zdefiniować jako bezgradientowe wibracje pola magnetycznego.) Odkryty wówczas mechanizm i zjawisko telepatii po raz pierwszy opublikowane zostało 9 stycznia 1996 roku w monografii [3] (patrz podrozdział D13 w [3]). W 1997 roku powtórzone one były w monografii [3/2]. Do niniejszej monografii trafiły poprzez monografię [1/3], w której opisano je w podrozdziale H13. Obecne ich opisy zawarte są w podrozdziale H7.1 niniejszej monografii.

#11. Narodziny idei "rezonatora magnetycznego". Odkrycie mechanizmu telepatii uświadomiło mi kolejny istotny atrybut komór oscylacyjnych. Atrybut ten odkryłem w wyniku swych badań nad telepatycznymi stacjami nadawczo-odbiorczymi. Takie stacje, czyli urządzenia które w sposób techniczny formują modulowane fale telepatyczne, muszą działać na zasadzie wykorzystania zjawiska "rezonowania" statycznych pól magnetycznych - patrz podrozdział H7.1. Ich najważniejszym więc pozespołem będzie urządzenie jakie można by nazywać "rezonanotorem magnetycznym". W zadaniu zbudowania pierwszego takiego urządzenia, jesteśmy więc obecnie w sytuacji dawnych naukowców eksperymentujących z polami elektrostatycznymi. Wszakże oni też nie mieli jeszcze pojęcia, że kiedyś odkryty zostanie elektryczny obwód drgający (czyli "rezonator elektryczny"). Dopiero ów obwód drgający wprowadził te pola w stan wibrowania. W ten sposób stworzył on fundamenty dla dzisiejszej radiokomunikacji, elektroniki i cybernetyki. Przez analogię do tamtej historycznej sytuacji, nasza obecna wiedza stałych pól magnetycznych jest jedynie na początku swej drogi do poznania. Spory więc czas zapewne jeszcze upłynie, zanim zbudowany zostanie nasz pierwszy działający "rezonator magnetyczny". Rezonator taki w przyszłości otworzy dla ludzi wykorzystanie wibracji pola magnetycznego dla różnorodnych celów technicznych, podobnie jak elektryczny obwód drgający otworzył kiedyś dla ludzi wykorzystanie wibracji pola elektrycznego.

Od początku swych prac nad komorą oscylacyjną intuicyjnie wiedziałem, że moje urządzenie reprezentuje właśnie pierwszą ideę takiego rezonatora magnetycznego. Dla wibracji magnetycznych idea ta już obecnie wprowadziła ten sam przełom poznawczy, co obwód oscylacyjny Henry'ego uczynił dla drgań elektrycznych. Pomimo jednak rozpracowania komory oscylacyjnej, oraz na przekór pełnego poznania jej zasady działania, budowy i najważniejszych atrybutów, nie było jeszcze jasne dla mnie do czego sprowadza się esencja idei rezonatorów magnetycznych. Esencja ta stała się oczywista dopiero po szczegółowym rozpracowaniu budowy i działania drugiego takiego rezonatora, oraz po odkryciu natury i mechanizmu fal telepatycznych. (Ten drugi rezonator przyjął postać "baterii telekinetycznej" opisanej w podrozdziale K2.4 niniejszej monografii.) Okazało się wtedy, że "esencją idei

rezonatorów magnetycznych jest iż urządzenia te reprezentują lustrzane odbicie elektrycznych obwodów drgających".

Aby wyjaśnić tutaj ową esencję, w sensie zasady swego działania rezonatory magnetyczne mimikują w sposób lustrzany zasadę działania elektrycznych obwodów drgających. Lustrzaność owego mimikowania polega w nich na symetrycznym odwracaniu tej zasady. (Jak pamiętamy, elektryczne obwody drgające już od dawna używane są do wytwarzania wibracji elektrycznych w urządzeniach elektronicznych, a także do formowania fal elektromagnetycznych w urządzeniach łączności telekomunikacyjnej.) Przykładowo, elektryczne obwody drgające muszą składać się z co najmniej dwóch podstawowych komponentów, tj. pojemności elektrycznej "C", oraz inercji magnetycznej "L" (zwanej też induktancją). Dlatego rezonatory magnetyczne muszą podobnie zawierać co najmniej dwie składowe, tj. inercję elektryczną "J" oraz pojemność magnetyczną "P". (Oczywiście, na dodatek do tych dwóch "lustrzanych" składowych, obie grupy urządzeń, tj. zarówno rezonatory magnetyczne jak i elektryczne obwody drgające, zawierały będą również oporność "R".) Podstawowe dedukcje przeprowadzone przeze mnie na ten temat opisane są w podrozdziałach K2.4 i N2.4 z niniejszej monografii. Ujawniają one, że wymaganej inercji elektrycznej "J" dostarczają rezonatorom magnetycznym świecące się (wzbudzone) jony, np. z mieszaniny rtęci i soli. Natomiast wymaganej pojemności magnetycznej "P" dostarcza im specjalnie ukształtowana przestrzeń odbijająca wibracje magnetyczne. Przestrzeń tą można nazwać "magnetyczną komorą rezonansową". (Przykładem najprostszej magnetycznej komory rezonansowej jest objętość czyli wnętrze powszechnie znanej piramidy.)

Narodziny i skryształizowanie się idei rezonatora magnetycznego otwierają drogę do technicznego urzeczywistnienia różnych wersji urządzeń telepatycznych. Umożliwiają one także teoretyczne sformułowanie wymogów ich działania, zrębów ich modeli matematycznych, itp. Ponadto ujawniają one, że nasza planeta jest naturalnym rezonatorem magnetycznym przechwytyjącym przenikające ją fale telepatyczne i przekazującym te fale do organizmów ludzi, zwierząt i roślin - patrz podrozdział D4 monografii [5/3].

#12. Wykorzystanie komory oscylacyjnej jako nadajnika i odbiornika telepatycznego. Po rozpracowaniu idei rezonatorów magnetycznych, uzyskałem pewność, że komora oscylacyjna jest właśnie jednym z nich. Komorę tą już od dawna uważałem za zdolną do wysyłania i odbierania sygnałów telepatycznych. Wszakże spełnia wszelkie wymagania aby pracować jako efektywny nadajnik i odbiorniki telepatyczny. W ten sposób idea rezonatorów magnetycznych wskazała dokładny sposób i zasadę na jakich komorę oscylacyjną można dodatkowo adoptować do funkcjonowania jako nadajnik i odbiornik telepatyczny o ogromnej mocy. Pierwsze opublikowanie idei działania komory oscylacyjnej jako nadajnika i odbiornika telepatycznego nastąpiło w monografii [3]. Potem powtórzone ono było w monografiach [3/2], [1/2], [1/3], a obecnie i w niniejszej monografii.

#13. Rozpracowanie warunków konstrukcyjnych komór oscylacyjnych drugiej i trzeciej generacji. Dokonane ono zostało już w trakcie opracowywania monografii [1/2]. W tamtej monografii położyłem bowiem szczególnie duży nacisk na uwypuklenie i dokładne opisanie różnic pomiędzy magnokraftami (i UFO) pierwszej, drugiej, oraz trzeciej generacji. Wszakże skoro nie potrafimy jeszcze ich budować, istotne jest abyśmy byli w stanie chociaż je rozpoznać, jeśli ktoś użyje je w naszej obecności. Wygląd zaś komór oscylacyjnych zabudowanych do pędników owych statków jest jedną z wizualnie najłatwiej odnotowalnych takich różnic. Aby więc uświadomić czytelnikom swoich monografii [1/2], a potem [1/3], jakie są różnice w wyglądzie komór oscylacyjnych poszczególnych generacji, a także różnice wyglądu wylotów z pędników magnetycznych wykorzystujących komory poszczególnych generacji, postanowiłem rozpracować i graficznie zilustrować ich wygląd. W tym jednak celu najpierw musiałem matematycznie rozwiązać zbiór warunków konstrukcyjnych jakim komory oscylacyjne poszczególnych generacji podlegają. Zbiór tych warunków opisany został w podrozdziałach C7.1.1 i C7.1.2, oraz C7.2.1 do C7.2.3, niniejszej monografii. Natomiast graficznie zilustrowany on został na rysunkach C8 i C11. Ich rozpracowanie nastąpiło w sierpniu 1997 roku, zaś upowszechniane one były w egzemplarzach monografii [1/2] jakie

rozsyłałem do Polski poczynawszy od września 1997 roku. Następnie były one upowszechniane we wszystkich egzemplarzach monografii [1/3].

#14. Rozpracowanie "prototypowych konfiguracji krzyżowych" formowanych z komór oscylacyjnych. Analizy zasady działania i technologii wykonania konfiguracji komór oscylacyjnych nazywanych "kapsuła dwukomorowa", ujawniły mi z czasem, że w pierwszych magnokraftach kapsuła taka nie będzie mogła jeszcze zastać użyta. Wszakże jest ona zbyt trudna do szybkiego technicznego wykonania. (Np. sygnały sterujące muszą w niej dostać się bezprzewodowo do "wewnętrznej" pływającej komory oscylacyjnej, przebijając się przez bardzo skoncentrowane pole magnetyczne komory "zewewnętrznej".) Dlatego też dla użytku pierwszych magnokraftów budowanych na Ziemi, rozpracowałem konstrukcje tzw. "prototypowych konfiguracji krzyżowych". Owe prototypowe konfiguracje krzyżowe są znacznie prostrze w budowie. Z powodzeniem mogą one więc zastępować w naszych pierwszych magnokraftach owe skomplikowane technologicznie kapsuły dwukomorowe. Prototypowe konfiguracje krzyżowe opublikowałem po raz pierwszy w 1998 roku w monografii [1/3]. W niniejszej monografii są one opisane w podrozdziale C7.2.1.

#15. Stopniowe zamarcie praktycznych prac wykonawczych i rozwojowych nad komorą oscylacyjną. Dla wynalazcy nie istnieje już chyba bardziej gorzki obraz, niż obserwować jak jego wynalazek stopniowo umiera. Niestety również i takie przykre doświadczenie stało się moim udziałem. Ogromnie intensywne w latach 1984 do 1988 prace badawczo rozwojowe nad komorą oscylacyjną, poczynawszy od około roku 1990 stopniowo zaczęły zamierać, aby całkowicie zaniknąć około 2000 roku. Około 2000 roku, ostatni znany mi badacz w Polsce, jaki ciągle zajmował się praktyczną budową komory oscylacyjnej zaprzestał swoich badań. W chwili obecnej (2004 rok) nad tym rewolucyjnym urządzeniem nikt już nie prowadzi praktycznych działań rozwojowych. Jest to ogromna strata dla naszej cywilizacji, bowiem komora oscylacyjna nosi w sobie potencjał aby otworzyć ludzkości wrota do gwiazd, a także aby kompletnie zrewolucjonizować wszystkie apseki naszego życia. Jedyne prace nad tą komorą, jakie ciągle są kontynuowane, to moje teoretyczne rozważania nastawione na udoskonalenie naszego zrozumienia jej budowy i działania, oraz na spopularyzowanie jej idei. Niemniej, dla powodów wyjaśnionych w podrozdziale A4, narazie nie zdołałem uzyskać dostępu do warunków jakie umożliwiałyby mi rozpoczęcie prac praktycznych nad zbudowaniem tego rewolucyjnego urządzenia.

* * *

Komentując powyższą historię wynalezienia i dotychczasowego rozwoju komory oscylacyjnej, powinienem tutaj podkreślić, że to urządzenie wynalezione zostało w rezultacie moich zawodowych zainteresowań naukowych wynikających z wykonywanej pracy wykładowcy uczelni technicznej. Podobnie też liczne inne urządzenia rozpracowane przeze mnie, jakich opisy zawarte są w moich monografiach, również wynalezione zostały w rezultacie moich zainteresowań zawodowych.

Podkreślając powyższe, jednocześnie muszę jednak natychmiast dodać, że wszystkie uczelnie w których kiedykolwiek pracowałem, nigdy nie aprobowaly moich badań nad tymi urządzeniami. W wielu przypadkach uczelnie te aktywnie nawet zwalczały owe badania i prześladowały mnie za ich prowadzenie. W rezultacie tego z biegiem czasu zmuszony zostałem do konspiracyjnego ukrywania przed przełożonymi i kolegami zawodowymi prawdziwego tematu swoich badań naukowych. Zmuszany też byłem do nieujawniania swoich osiągnięć twórczych w przypadkach poszukiwania nowej pracy, a także do podejmowania prac leżących znacząco poniżej mojego poziomu doświadczeń i zdolności twórczych. Z kolei fakt zmuszania naukowca do prowadzenia swoich badań w konspiracji, świadczy że sytuacja na uczelniach w latach początku 21 wieku osiągnęła poziom niemal parodii celów jakim nauka i naukowcy powinni służyć. Narasta więc paląca już konieczność zreformowania naszych instytucji naukowych, tak jak po mrokach i prześladowaniach średniowiecza koniecznym było zreformowanie instytucji kościoła (więcej na temat owej reformacji dzisiejszej nauki wyjaśnione jest w podrozdziale H10). Nienaruszalnym przy tym fundamentem takiego reformowania, powinno być prawo naukowców totalizacyjnych do wolności w myśleniu. Prawo

to powinno zostać nie tylko wyraźnie i jednoznacznie zadeklarowane w dokumentach konstytucyjnych uczelni, ale także powinno być respektowane w praktyce. Prawo to powinno stwierdzać, że "naukowiec może być rozliczany tylko z tego jak produktywnie i twórczo pracuje, nigdy zaś z tego nad czym pracuje. Wszakże każda tematyka badań w efekcie końcowym zawsze służy społeczeństwu i ludzkości". Jak wiadomo, dotychczas takie prawo do wolności myślenia nie jest jeszcze praktycznie urzeczywistniane, czy nawet oficjalnie deklarowane w dokumentach konstytucyjnych uczelni. Stąd naukowcy zamiast "poszukiwać i ujawniać prawdę bez względu na to jaka by ona nie była", raczej dla zabezpieczenia sobie braku kłopotów typowo "zatajają i ignorują tę prawdę która nie odpowiada im samym lub komuś ważnemu, niebezpiecznemu, czy krzykliwemu" - patrz punkty #11 i §11 w podrozdziale JB6. Z kolei wynalazcy i odkrywcy którzy faktycznie są jedynymi którzy dokładają nowy wkład do rozwoju naszej cywilizacji, zamiast być za to nagradzani, w rzeczywistości są karani i praktycznie zamieniani w męczenników.

Wszystkie odkrycia i wynalazki opisane w tej monografii zostały zapracowane ciężkim wysiłkiem i żmudnymi badaniami. Wymagały one długotrwałych przemyśleń, weryfikacji, modyfikacji i usprawnień. Z kolei ich sformułowanie końcowe wyniknęło z poziomu mojego poznania obecnej nauki i techniki ziemskiej osiągniętego na drodze długotrwałego i żmudnego studiowania. W moim więc przypadku, wynalazki te niestety nie pojawiły się od razu w gotowej postaci, zainspirowane bez wysiłku w jakiś cudowny sposób czy przekazane mi przez nadrzędne istoty. Miejmy więc nadzieję, że jako takie uznane one zostaną za "zapracowane" wystarczająco wysokim trudem i wyrzeczeniami, aby włączone kiedyś mogły zostać do trwałego dorobku technicznego naszej cywilizacji. (Po wyjaśnienie czym jest owo "zapracowanie", patrz Prawo Moralne Zapracowania na Wszystko wyjaśnione w punkcie #3A podrozdziału I4.1.1 niniejszej monografii.)

C3. Zasada działania komory oscylacyjnej

Prąd elektryczny przepływający przez zwoje przewodnika nie jest jedynym źródłem kontrolowanego pola magnetycznego. Innym dobrze znanym takim źródłem jest iskra elektryczna, tj. zjawisko reprezentujące energię elektryczną w najczystszej z postaci. Znanych jest wiele sposobów na wytwarzanie iskier elektrycznych, jednakże dla zastosowań opisywanych tutaj najprzydatniejszym jest tzw. "obwód oscylacyjny z iskrownikiem". Unikalną własnością takiego obwodu jest jego zdolność do pochłaniania, sumowania i najefektywniejszego wykorzystywania energii elektrycznej dostarczonej do niego. Energia ta następnie manifestuje się w postaci zwolna zanikającej serii iskier elektrycznych wytwarzanych przez ten obwód.

Wynalezienie obwodu oscylacyjnego z iskrownikiem nastąpiło w roku 1845 przez fizyka amerykańskiego, Joseph'a Henry. Zauważył on, że jeśli rozładować butelkę lejdejską (Layden jar) poprzez uzwojenia induktora, wtedy otrzymywało się oscylującą iskrę. W kilka lat potem Lord Kelvin, fizyk i inżynier angielski, dowiódł matematycznie że wyładowanie w tak skonstruowanym obwodzie musi następować w sposób oscylacyjny.

W tym miejscu należy podkreślić że obwód oscylacyjny z iskrownikiem był pierwszym obwodem wynalezionym na naszej planecie jaki wytwarzał drgania elektryczne. Jego zbudowanie wnosilo więc równie rewolucyjne konsekwencje jak przykładowo opracowanie pierwszej maszyny parowej. Obwód ten dostarczył bowiem fundamentów poznawczych dla późniejszego uformowania wielu oddzielnych dziedzin nauki bazujących właśnie na drganiach elektrycznych, przykładowo elektroniki czy cybernetyki, zaś na jego zasadzie oparte jest działanie ogromnej ilości dzisiejszych urządzeń, takich jak radia, telewizory, komputery, przyrządy pomiarowe, oraz wiele innych. Powinniśmy więc pamiętać i szczerze potwierdzać, że gdyby nie odkrycie Henry'ego wówczas nasza cywilizacja nie byłaby na poziomie na jakim znajduje się obecnie.

C3.1. Inercja elektryczna induktora stanowi siłę motoryczną dla oscylacji w tradycyjnym obwodzie oscylacyjnym z iskrownikiem

Rysunek C1 (a) pokazuje tradycyjną konfigurację obwodu elektrycznego z iskrownikiem, tzn. konfigurację wynalezioną przez Henry'ego. Najbardziej wyróżniającą się cechą tego obwodu jest, iż powstaje on poprzez połączenie razem w jeden obwód zamknięty trzech istotnych elementów, tj. L , C_1 i E , jakie przyjmują formę oddzielnych części lub urządzeń. Części te to: (1) induktor L , zawierający długi przewód zawinięty w wiele zwojów, który dostarcza obwodowi własności zwanej "indukcyjność"; (2) kondensator C_1 , którego własność zwana "pojemnością elektryczną" umożliwia obwodowi gromadzenie ładunków elektrycznych; (3) iskrownik E , którego dwie równoległe elektrody, prawa E_R i lewa E_L oddzielone od siebie warstewką gazu, wprowadzają "przerwę iskrową" do obwodu.

Kiedy na płytki P_F i P_B kondensatora C_1 przyłożone zostają ładunki elektryczne "+q" i "-q", powoduje to przepływ prądu elektrycznego "i" poprzez przerwę iskrową E oraz induktor L . Prąd "i" musi przejawiać się w postaci iskier "S" i musi także wytworzyć strumień magnetyczny "F". Mechanizm kolejnych transformacji energii następujących w induktorze L (jaki niezależnie od niniejszego podrozdziału opisany jest także w wielu książkach z zakresu elektroniki i fizyki) powoduje, iż iskra "S", gdy już raz się pojawi pomiędzy elektrodami E , będzie oscylowała tam tak długo, aż energia obwodu zostanie rozproszona.

Obwód oscylacyjny z iskrownikiem reprezentuje elektryczną wersję wielu istniejących obecnie urządzeń jakie wytwarzają jedno z najbardziej powszechnych w naturze zjawisk, tj. ruch drgający. Analogią mechaniczną do tego obwodu, znaną doskonale każdemu, jest huśtawka. We wszystkich urządzeniach wytwarzających taki ruch, tj. zarówno w obwodzie oscylacyjnym jak i huśtawce, pojawienie się oscylacji wywoływane jest działaniem Zasady Zachowania Energii. Zasada ta powoduje, iż energia początkowa dostarczona do takiego urządzenia oscylacyjnego, zostaje następnie w nim uwięziona w procesie nieustannie powtarzających się transformacji w dwie formy: energii potencjalnej i energii kinetycznej. W przypadku obwodu oscylacyjnego z iskrownikiem, "energia potencjalna" reprezentowana jest przez pole elektryczne przeciwnych ładunków elektrycznych "+q" i "-q" zgromadzonych na obu okładzinach kondensatora - patrz rysunek C1 (a). Właśnie różnica potencjałów elektrycznych spowodowana obecnością tych ładunków, formuje siłę motoryczną jaka wymusza przepływ prądu "i" poprzez dany obwód. W przypadku huśtawki, ta sama energia potencjalna zostaje wprowadzona na drodze odchylenia jej ramienia, wraz z zamocowanym do niego siedzeniem, od położenia pionowego. W rezultacie, ciężar z danej huśtawki (np. siedzące na niej dziecko) wzniesiony zostanie na określoną wysokość. Energia potencjalna tego ciężaru wymusza później jego przyspieszanie w dół do pozycji równowagi, transformując się w ten sposób stopniowo w energię kinetyczną. W dolnym punkcie huśtawki cała energia potencjalna przetransformowana już zostaje na energię kinetyczną, która manifestuje się w postaci szybkiego ruchu ciężaru przyłożonego do jej siedzenia. W obwodzie oscylacyjnym z iskrownikiem druga z form energii, tj. energia kinetyczna, manifestuje się w formie strumienia "F" pola magnetycznego wytwarzanego przez induktor L .

Wzajemne transformowanie się energii potencjalnej w energię kinetyczną, i vice versa, wymaga istnienia rodzaju pośrednika jaki aktywuje mechanizm zamiany tej energii. Ów pośrednik jest wprowadzany przez element odznaczający się własnością zwaną "inercja". Inercja jest więc jednym z najistotniejszych czynników napędowych podtrzymujących oscylacje w dowolnym układzie drgającym. Działa ona jak rodzaj "pompy" jaka wymusza transformacje energii z formy potencjalnej, w formę kinetyczną, a potem znowóż w odwrotną formę potencjalną. Owa "pompa" zawsze stara się odtworzyć początkową ilość energii potencjalnej, istniejącą w chwili rozpoczęcia się danego cyklu oscylacji, pomniejszoną jedynie przez wielkość jej rozproszenia następującego w czasie oscylacji. Stąd też element inercyjny jest najbardziej istotnym składnikiem każdego układu drgającego. W obwodzie oscylacyjnym z

iskrownikiem jego funkcję wypełnia induktor L, którego indukcyjność (wyrażona w jednostkach zwanych "henry") reprezentuje inercję elektryczną. W huśtawce, inercja mechaniczna jest dostarczana przez masę jej ładunku (wyrażoną w kilogramach). Powyższe jest powodem dla którego indukcyjność w oscylacjach elektrycznych jest uważana za odpowiednik masy w drganiach mechanicznych.

Aby zwiększyć inercję mechaniczną koniecznym jest dołożenie dodatkowej masy do tej już uczestniczącej w transformacjach energii. Natomiast zwiększenie inercji elektrycznej wymaga wydłużenia przewodnika w którym przepływający przez niego prąd elektryczny wystawiony zostaje na działanie swego własnego pola magnetycznego. Praktycznie uzyskiwane jest to poprzez budowanie induktorów zawierających większą liczbę zwojów tego samego przewodnika, zawiniętych ciasno jeden przy drugim, tak aby każdy z nich znajdował się w zasięgu pola magnetycznego produkowanego przez inne zwoje.

Przeanalizujmy teraz mechanizm wytwarzania drgań elektrycznych w obwodzie oscylacyjnym z iskrownikiem pokazanym na rysunku C1 (a). Załóżmy, że w początkowej chwili czasowej $t=0$ elektrody PB i PF kondensatora C1 posiadają zgromadzone już na nich przeciwstawne ładunki "+q" i "-q", oraz że prąd elektryczny "i" w induktorze L wynosi zero. W owej chwili czasowej cała energia obwodu jest więc przechowywana na okładzinach kondensatora C1 w formie energii potencjalnej. Przeciwstawne ładunki zgromadzone na okładzinach kondensatora C1 wywołują siłę elektromotoryczną jaka inicjuje przepływ prądu "i". Aby ułatwić interpretację zachowania się iskier, w tej publikacji **prąd elektryczny zdefiniowany jest jako ruch elektronów od bieguna negatywnego do pozytywnego**. Prąd "i" pojawia się na elektrodach E jako iskra elektryczna "S", podczas gdy w induktorze L wytwarza on strumień magnetyczny "F". Podczas gdy różnica ładunków "q" zakumulowanych na okładzinach kondensatora C1 ulega zmniejszeniu, ich energia potencjalna mająca formę pola elektrycznego także ulega zmniejszeniu. Owa energia zostaje przetransformowana w pole magnetyczne jakie pojawia się wokół induktora w efekcie przepływu przez niego prądu "i". Stąd w tej pierwszej fazie oscylacji, jaką możemy nazwać fazą **aktywną**, pole elektryczne się zmniejsza, pole magnetyczne się zwiększa, zaś energia zostaje transformowana z postaci potencjalnej w postać kinetyczną płynąc od kondensatora C1 do induktora L. Kiedy cały ładunek na kondensatorze C1 zaniknie, pole elektryczne w tym kondensatorze osiągnie zero, zaś energia potencjalna przechowywana w tym polu w całości zostanie przetransformowana w pole magnetyczne "F" induktora L. Siła elektromotoryczna jaka poprzednio powodowała przepływ prądu "i" zostaje więc wyeliminowana. Jednakże prąd w przewodniku kontynuuje transportowanie negatywnych ładunków z płyty PB kondensatora C1 do płyty PF, właśnie z powodu działania inercji elektrycznej. Inercja ta powstrzymuje prąd "i" (a więc także iskrę "S") przed zaniknięciem i utrzymuje jego przepływ kosztem energii kinetycznej zawartej w polu magnetycznym. W tej drugiej więc fazie oscylacji, jaką możemy nazwać fazą **inercyjną**, energia przepływa już z induktora L z powrotem do kondensatora C1, powodując ponowne odbudowywanie istniejącego tam uprzednio pola elektrycznego. Stopniowo cała energia kinetyczna pola magnetycznego zostaje przetransformowana z powrotem na energię potencjalną kondensatora C1. Po owym odbudowaniu, sytuacja uzyskana w owej chwili $t=(1/2)T$ podobna będzie do sytuacji początkowej w chwili $t=0$, tyle tylko że kondensator będzie teraz naładowany w odwrotny sposób. W dalszej fazie oscylacji kondensator C1 rozpocznie ponowne rozładowanie, zaś cały opisany powyżej proces powtórzy się ponownie, tym razem w przeciwnym kierunku. Po czasie $t=T$ (gdzie "T" jest tzw. "okresem pulsacji" lub "okresem oscylacji" danego obwodu) sytuacja wróci więc do stanu identycznego jak w chwili początkowej $t=0$. Raz więc zastartowane, takie oscylacje będą trwały aż oporność procesu rozproszy energię wymaganą do jego podtrzymywania.

C3.2. W zmodyfikowanym obwodzie oscylacyjnym z iskrownikiem inercji elektrycznej dostarczy indukcyjność iskry elektrycznej

Wiadomo że iskry elektryczne są nośnikiem bardzo wysokiej inercji elektrycznej. Stąd iskry te posiadają zdolność zastąpienia zwojów induktora w dostarczeniu obwodowi oscylacyjnemu wymaganej indukcyjności. Istnieją jednak dwa warunki tego zastąpienia, tj.: (1) iskra musi posiadać odpowiednią długość aktywną, oraz (2) droga iskry musi przebiegać w zasięgu wytwarzanego przez siebie pola magnetycznego. Aby wypełnić obydwie te warunki, niemożliwe jest powtórzenie rozwiązania konstrukcyjnego użytego w induktorze, z prostej przyczyny iż iskra elektryczna będzie opierała się naszym próbom zawijania jej w kilka kolejnych zwojów. Jednakże ten sam efekt może zostać osiągnięty w odmienny sposób. **Wymaganej indukcyjności jest też w stanie dostarczyć cały snop isker przeskakujących równocześnie po równoległych trajektoriach, każda z których zastępuje akcję pojedynczego zwoju induktora.** Indywidualne iskry w takim snopie będą więc odpowiednikami poszczególnych zwojów induktora. Stąd jeśli ilość isker osiągnie wymaganą liczbę, wszystkie razem będą one w stanie dostarczyć obwodowi wymaganej indukcyjności.

Rysunek C1 (b) ukazuje wersję typowego obwodu oscylacyjnego z iskrownikiem, jaką celowo zmodyfikowałem, a jaka właśnie wykorzystuje do swego działania inercję snopa równoległych isker. Najbardziej wyróżniającą się cechą tej wersji jest, iż wszystkie trzy niezbędne składniki obwodu Henry'ego, tj. indukcyjność L , pojemność $C1$, oraz przerwa iskrowa E , są w niej dostarczane przez pojedyncze urządzenie w postaci pary elektrod. Stąd to jedno urządzenie zastępuje wszystkie trzy składniki tradycyjnego obwodu. Tem mój zmodyfikowany obwód oscylacyjny z iskrownikiem składa się więc z owej pary przewodzących elektrod P_F i P_B , jakie umocowane zostały do dwóch przeciwstawnych ścian komory sześcienniej wykonanej z materiału izolacyjnego (np. szkła) i wypełnionej gazem dielektrycznym. Każda z tych elektrod podzielona została na wiele małych segmentów odizolowanych nawzajem od siebie. W części (b) rysunku C1 segmenty te oznaczono numerami 1, 2, 3, ..., p . Każda para segmentów ustawionych naprzeciwko siebie tworzy pojedynczy elementarny kondensator. Na rysunku C1 (b) każda taka para segmentów formujących elementarny kondensator oznaczona została tym samym numerem, np. "3" lub "p". Z kolei ów kondensator, po otrzymaniu odpowiedniego ładunku elektrycznego, przekształca się w parę elektrod wymieniających z sobą pojedynczą iskrę elektryczną (np. " S_3 " lub " S_p "). Stąd obie elektrody P_F i P_B omawianego obwodu wytwarzają tyle isker elektrycznych na ile segmentów zostały one podzielone. Suma owych isker przeskakujących w tym samym momencie w formie upakowanego snopa (pęku), dostarcza obwodowi wymaganej indukcyjności elektrycznej.

Podsumujmy teraz istotę modyfikacji obwodu Henry'ego jakiej dokonałem i wyjaśnię powyżej. Trzy oddzielne części/elementy składowe tradycyjnego obwodu oscylacyjnego (tj. induktor, kondensator i iskrownik), z których każdy wypełniał jedną funkcję, zastąpiono jedną częścią, która za to wypełnia aż trzy funkcje równocześnie. Ta jedna część/element to para przewodzących elektrod zamocowanych do dwóch przeciwstawnych ścianek komory sześcienniej i podzielonych na małe segmenciki. Każdy z tych indywidualnych segmencików wypełnia funkcję elementarnego kondensatora i iskrownika, zaś formowany przez nie wszystkie pęk równoległych isker wypełnia też funkcję induktora.

Jeśli rozpatrzeć zasadę działania zmodyfikowanego w ten sposób obwodu oscylacyjnego, staje się oczywiste, że zasada ta jest identyczna jak w przypadku obwodu Henry'ego. Po rozpoczęciu równomiernego nasycania ładunkami elektrycznymi wszystkich segmencików obu elektrod P_F i P_B , energia potencjalna obwodu zaczyna się zwiększać. Gdy różnica potencjałów pomiędzy elektrodami przekroczy wartość przebicia " U ", wyładowanie iskrowe zostaje rozpoczęte. Wyładowanie to przyjmuje formę pęku równoległych isker $S_1, S_2, S_3, \dots, S_p$, przeskakujących pomiędzy położonymi naprzeciwko siebie segmencikami obu elektrod. W pierwszej więc, **aktywnej** fazie cyklu oscylacyjnego pole magnetyczne formowane przez każdą z tych isker będzie stopniowo absorbowало energię zgromadzoną początkowo w formie potencjalnej na obu elektrodach. Kiedy potencjały obu elektrod P_F i P_B się zrównają, inercja elektryczna isker będzie kontynuowała przepompowywanie ładunków pomiędzy nimi, transformując energię kinetyczną zawartą w polu magnetycznym z powrotem w energię

potencjalną pola elektrycznego. Stąd na końcu drugiej, **inercyjnej** fazy oscylacji iskiek obie elektrody znowu będą zawierały początkowe ładunki, tyle tylko że odwrotnego znaku. Następnie cały proces zostanie powtórzony w odwrotnym kierunku. Jeśli więc niewielkie rozpraszanie energii zachodzące podczas tego procesu zostanie jakoś uzupełnione, opisany tu proces może powtarzać się bez końca.

Takie działanie zmodyfikowanego obwodu oscylacyjnego uwalnia wszystkie zachodzące w nim zjawiska z więzów materiałowych. W rezultacie więc prąd elektryczny nie musi przepływać przez przewodnik zaś jego natężenie nie jest ograniczane własnościami przewodzącymi materiałów użytych na zwoje. Ponadto zjawiska elektryczne zostają wyeksponowane na działania sterujące jakie umożliwiają ich ukształtowanie w pożądanym kierunku. Są to bardzo istotne osiągnięcia i jak to zostanie wykazane potem, stanowią one źródło wielu zalet operacyjnych opisywanego tu urządzenia.

Pęki iskiek oscylujących w urządzeniu pokazanym na rysunku C1 (b) produkowały będą zmienne pole magnetyczne. Ponieważ snopy te przeskakują po tej samej drodze w obu kierunkach, pole to będzie nosiło charakter wirowy ("vortex") podobny do pola otaczającego odcinek prostego przewodnika (tj. wszystkie jego linie sił będą formowały obiegi we wzajemnie do siebie równoległych płaszczyznach). Takie pole nie będzie więc wykazywało jednoznacznej biegunowości, ponieważ jego bieguny magnetyczne N i S nie posiadają umiejscowienia. Stąd aby wytworzyć dwubiegunowe pole magnetyczne ze stałą pozycją jego biegunów N i S, koniecznym jest kontynuowanie rozwoju omawianej tu zmodyfikowanej wersji obwodu oscylacyjnego o jeden dalszy krok.

C3.3. Zestawienie razem dwóch zmodyfikowanych obwodów formuje komorę oscylacyjną wytwarzającą dipolarne pole magnetyczne

Końcową postać omawianego tutaj obwodu pokazano na rysunku C1 (c). Jest to właśnie postać jakiej nadano nazwę "komora oscylacyjna". Komora oscylacyjna zostaje uzyskana poprzez złożenie z sobą dwóch obwodów na rysunku C1 (c) oznaczonych jako C_1 and C_2 . Każdy z tych obwodów jest identyczny do tego omówionego w poprzednim podrozdziale i pokazanego w części (b) rysunku C1. Stąd komora taka składa się z czterech posegmentowanych elektrod, oznaczonych jako P_F , P_B , P_R i P_L , tj. przedniej (po angielsku "front"), tylnej ("back"), prawej ("right") i lewej ("left"). Każda z tych elektrod również podzielona została na taką samą liczbę "p" segmencików oraz ustawiona jest naprzeciwko identycznej elektrody formując z nią razem jeden z obu nawzajem kooperujących obwodów. Oba te obwody produkują cztery pęki iskiek na rysunku C1 (c) oznaczonych jako S_{R-L} , S_{F-B} , S_{L-R} , i S_{B-F} , które przeskakują pomiędzy przeciwległymi elektrodami. Pęki te pojawiają się w ściśle zdefiniowanej kolejności, jeden po drugim, posiadając wzajemne przesunięcie fazowe pomiędzy kolejnymi przeskokami wynoszące jedną czwartą ($1/4$) okresu T ich całkowitej sekwencji przeskoków (tj. " $(1/4)T$ ").

Zanim mechanizm wyładowań iskrowych w tej końcowej konfiguracji komory zostanie opisany, powinniśmy przypomnieć sobie tutaj działanie elektromagnetycznych sił odchylających jakie będą starały się wyprzeć iskry z zasięgu dipolarnego pola magnetycznego. Siły te są tymi samymi siłami jakie powodują eksplozję zwojów w potężnych elektromagnesach (ich omówienia już dokonano w punkcie #1 podrozdziału C1). W przypadku oscylujących iskiek siły te będą przypierały ich pęki do lewej elektrody wzdłuż powierzchni której dane wyładowanie następuje. Dla przykładu wszystkie iskry w pęku S_{R-L} skaczące od płyty P_R do płyty P_L będą dociskane do powierzchni płyty P_F (w owym momencie płyta P_F zwiększa swój ujemny ładunek). Z tego powodu poszczególne iskry formujące kolejne pęki S_{R-L} , S_{F-B} , S_{L-R} , i S_{B-F} , zamiast przecinać nawzajem drogi innych iskiek, będą uginały własną drogę do powierzchni ścianek komory leżących po ich lewej stronie, wytwarzając w ten sposób rodzaj zgodnie obiegającego, uporządkowanego wiru iskrowego. Warto zauważyć, że płyta wzdłuż powierzchni której iskry w danym momencie czasowym przeskakują, jest

zabezpieczona przed ich wniknięciem do jej materiału. Owo zabezpieczenie wynika głównie z uformowania elektrod z dużej liczby małych segmencików (igieł), każdy z których jest odizolowany od innych podobnych segmencików. Stąd oporność dla przepływu prądu w poprzek elektrody nie jest mniejsza od oporności wyładowania poprzez dielektryczny gaz komory.

Założmy przez chwilę, że początkowe naładowanie komory oscylacyjnej jest dokonane w ten sposób, iż w chwili $t=0$ jako pierwszy pojawi się pęk isker oznaczony jako S_{R-L} , zaś po upływie okresu czasu równego $t=(1/4)T$ - pojawi się pęk S_{F-B} (porównaj rysunek C1 "c" z rysunkiem C4 "a"). Założmy także, iż od samego początku tych wyładowań, wzdłuż osi magnetycznej "m" komory panuje produkowany przez to urządzenie strumień magnetyczny "F". Strumień ten odchyła wszystkie iskry dociskając je do ich lewostronnych ścianek komory. Po początkowym naładowaniu kondensatora C2, w chwili czasowej $t=0$, pojawi się pęk isker oznaczony jako S_{R-L} , jaki przeskoczy od elektrody P_R do elektrody P_L . Iskry te wytworzą swoje własne pole magnetyczne o natężeniu " ΔF " wydatek jakiego doda się do całkowitego pola "F" już panującego w tej komorze. Pole "F" zagina drogę wszystkich przeskakujących isker, przypierając je do powierzchni elektrody P_F . W chwili czasowej $t = (1/4)T$ potencjały elektrod P_R i P_L wyrównują się, jednakże inercja elektryczna pęku isker S_{R-L} ciągle kontynuuje transportowanie ładunków od elektrody P_R do elektrody P_L , kosztem energii kinetycznej zakumulowanej w polu magnetycznym. W tym samym momencie czasowym $t = (1/4)T$ rozpoczyna się działanie drugiego obwodu oscylacyjnego, stąd zainicjowany zostaje przeskoczek pęku isker oznaczonego jako S_{F-B} . Podobnie jak pęk poprzedni, również i ten pęk wytwarza swój strumień magnetyczny " ΔF " jaki dodaje się do całkowitego strumienia "F" komory powodując m.in. przypieranie isker S_{F-B} do powierzchni elektrody P_L . Stąd w przedziale czasu od $t = (1/4)T$ do $t = (2/4)T = (1/2)T$, dwa pęki isker, S_{R-L} i S_{F-B} , współistnieją w komorze równocześnie. Pierwszy z nich - inercyjny, przetransformowuje energię z pola magnetycznego do pola elektrycznego, natomiast drugi - aktywny, transformuje energię pola elektrycznego w pole magnetyczne. W chwili czasowej $t = (2/4)T = (1/2)T$ elektrody P_L i P_R osiągają różnicę potencjałów równą różnicy początkowej (tj. w chwili $t=0$), jednak ich ładunki są teraz przeciwne niż początkowo. Stąd pęk isker S_{R-L} zaniknie, podczas gdy zainicjowany zostaje pęk S_{L-R} skaczący z kierunku do nich przeciwnym. Pęk ten przypierany jest do powierzchni elektrody P_B przez pole "F". W tej samej chwili czasowej $t = (2/4)T = (1/2)T$ elektrody P_F i P_B osiągają stan zrównania się ich potencjałów, stąd pęk isker S_{F-B} przechodzi w swoją inercyjną fazę. W przedziale czasu od $t = (2/4)T = (1/2)T$ do $t=(3/4)T$ w komorze znowóż współistnieją aż dwa pęki isker, tj. S_{F-B} i S_{L-R} , pierwszy z których - inercyjny, konsumuje pole magnetyczne, podczas gdy drugi - aktywny, wytwarza je. W chwili czasowej $t=(3/4)T$ iskry S_{F-B} zanikają zaś iskry S_{B-F} zostają wytworzone (przypierane do elektrody P_R), podczas gdy iskry S_{L-R} przechodzą do swojej inercyjnej fazy. W momencie czasowym $t = (4/4)T = (1)T$ iskry S_{L-R} również zanikają zaś iskry S_{R-L} zostają wytworzone (przypierane do elektrody P_F), podczas gdy iskry S_{B-F} przechodzą do ich inercyjnej fazy. W tym momencie więc cały cykl przeskoczków isker zostaje zamknięty, zaś sytuacja w czasie $t = (4/4)T = (1)T$ jest identyczna do sytuacji w chwili początkowej $t=0$. Stąd proces przeskoczków jaki nastąpi potem będzie już powtórzeniem procesu właśnie tu opisanego.

Powyższa analiza kolejności pojawień się oraz drogi pęków isker w komorze oscylacyjnej ujawnia bardzo pożądaną regularność. Owe pęki isker formują bowiem rodzaj wirującego ciągłego łuku elektrycznego, jakiego kompletny obieg wokół komory złożony zostaje z czterech nakładających się na siebie segmentów. Łuk ten obiega wokół osi magnetycznej komory zawsze w tym samym kierunku. W wyniku tego procesu, zgodnie z teorią elektromagnetyzmu, łuk ten musi produkować silne, pulsujące, dipolarne pole magnetyczne. Uzyskanie takiego pola koronuje więc długą i trudną drogę w moich poszukiwaniach nowej zasady wytwarzania pól magnetycznych, jaka eliminowałaby wady obecnie używanych w tym celu elektromagnesów.

C3.4. Igłowe elektrody

Konstrukcja komory oscylacyjnej opisana powyżej, stanowi pierwszy opis tej konstrukcji jaki kiedykolwiek opublikowałem. Jednak następujące później prace rozwojowe nad tą komorą wykazały, że konstrukcja ta jest trudna do zrealizowania z uwagi na początkowo sugerowany płytkowy kształt elektrod. Jak to bowiem już wspomniano w punkcie #7 podrozdziału C2, płytkowe elektrody sprzyjały przepływowi prądu iskier "na skróty" zamiast jak powinien on płynąć zgodnie z zasadą działania komory - patrz **rysunek C2**.

W toku dalszych badań eksperymentalnych udało się jednak ustalić, że zastosowanie elektrod igłowych zamiast płytkowych eliminuje ten problem - patrz część "b" rysunku C2. Stąd też w dalszej części tego rozdziału przez elektrody komory należy rozumieć igły wystające ku wewnątrz z jej ścianek i wypełniające wszystkie funkcje jakie przy objaśnianiu zasady jej działania nałożono na płaskie segmenty. (Mimo wprowadzenia elektrod igłowych, dla uproszczenia rozważań w objaśnieniach z początkowej części tego rozdziału ciągle zachowano płaskie segmenty elektrod. Wszakże formują one w umyśle czytelnika bardziej ilustracyjny system pojęciowy, bazujący na tradycyjnym zrozumieniu akumulatorów elektryczności jako dwóch płyt równoległych do siebie.)

C4. Przyszły wygląd komory oscylacyjnej

Nie jest trudnym zaspokojenie wymagań komory oscylacyjnej na materiały konstrukcyjne. Urządzenie to może być bowiem wykonane praktycznie z dowolnego materiału, zakładając że jego obudowa jest dobrym izolatorem elektryczności zaś jego elektrody zostały wykonane z dobrych przewodników elektryczności. Oba te materiały muszą też być magnetycznie obojętne, w przypadku bowiem użycia np. stali zostałyby one zniszczone wytwarzanym przez komorę polem magnetycznym. Stąd nawet tak starożytne materiały dostępne już przed tysiącami lat, jak drewno i złoto, wystarczają dla jej zbudowania. Jeśli przypadkiem zbudowana z tych pradawnych materiałów, komora oscylacyjna wyglądałaby jak niepozorna skrzynka czy kostka drewniana. Jej wygląd niczym nie zapowiadałby niezwyklej mocy ukrytej w jej wnętrzu.

Na naszym poziomie rozwoju dostępne są przezroczyste materiały izolacyjne, które również posiadają dużą wytrzymałość mechaniczną oraz są magnetycznie obojętne. Jednym z najpowszechniej występujących ich przykładów jest zwykłe szkło czy pleksiglas. Jeśli więc obudowę (ścianki) komory oscylacyjnej zbudować z takich właśnie przezroczystych izolatorów, wtedy użytkownik mógłby obserwować procesy zachodzące w jej wnętrzu, np. przeskoiki iskier elektrycznych, gęstość energii ciągle zawartej w komorze, działanie sterowania, itp. Współczesna elektronika wytworzyła również zapotrzebowanie na przezroczyste przewodniki. Już obecnie takie przewodniki można spotkać w niektórych zegarkach elektronicznych i kalkulatorkach. Jakość tych przezroczystych przewodników z czasem będzie ulegała poprawie, wkrótce więc prawdopodobnie możemy się spodziewać, iż ich własności elektryczne będą porównywalne do tych z dzisiejszych metali. Załóżmy więc, że w chwili zbudowania pierwszych działających komór oscylacyjnych ich budowniczy będą już w stanie wykonać je w całości z owych przezroczystych materiałów (tj. zarówno izolatorów jak i przewodników). Stąd zaciekawiony obserwator działania takich komór zobaczyłby przed sobą typowy "kryształ", tj. lśniąca kostkę sześcienną całą wyszlifowaną z przezroczystego materiału - patrz **rysunek C3**. Wzdłuż wewnętrznych powierzchni tej kryształowej kostki, jasno-żółte oscylujące iskry będą migotały. Iskry te sprawią wrażenie zamrożonych w tych samych pozycjach, aczkolwiek od czasu do czasu dokonujących nagłych poruszeń jak kłębowisko uśpionych ognistych węży. Ich drogi będą ciasno przylegały do wewnętrznych powierzchni ścianek komory, dociskane do nich przez elektromagnetyczne siły odchylające omówione już poprzednio. Wnętrze kostki będzie wypełnione potężnym pulsującym polem magnetycznym oraz rozrzedzonym gazem dielektrycznym. Pole to, gdy obserwowane z kierunku

prostopadłego do jego linii sił, będzie pochłaniało światło. Stąd sprawi ono wrażenie gęstego czarnego dymu wypełniającego wnętrze tego przeźroczystego kryształu.

Jest łatwe do zauważenia, że iskry elektryczne posiadają jakąś magiczną moc nad ludźmi. Kiedy na wystawie naukowej, albo podczas "dni otwartych" w laboratoriach, demonstrator uruchomi którąś z maszyn wytwarzających iskry, przykładowo cewkę Tesli, cewkę indukcyjną, lub maszynę Van de Graaff'a, widzowie nieodparcie przyciągani są do tego pokazu (tj. niemal "grawitują" do niego). Trzaski wyładowań i błyski iskieł zawsze posiadały jakąś tajemniczą, hipnotyczną moc jaka działa na każdego i jaka dostarcza niezapomnianych wrażeń. Potęga emanująca z wnętrza komory oscylacyjnej podobnie będzie przykuwała uwagę i wyobraźnię ludzi patrzących na jej działanie. Przyszli obserwatorzy tego urządzenia będą mieli odczucie patrzenia bardziej na jakieś żyjące stworzenie, zajęte wykonywaniem swoich fascynujących i tajemniczych czynności życiowych, niż na kawałek maszyny zajętej zwykłym procesem swego działania. Ogrom energii złapanej, okiełznanej, i przyczajonej we wnętrzu komory oscylacyjnej będzie fascynował widzów, pozostawiając ich z szeroką gamą żywych odczuć, wpisanych na zawsze do ich pamięci.

Obserwując ten niepozorny przeźroczysty kryształ, osoba patrząca będzie prawdopodobnie miała trudności z wyobrażeniem sobie, iż aby osiągnąć moment swojego narodzenia, owo urządzenie, tak przecież proste w kształtach, wymagało gromadzenia ludzkiej wiedzy i doświadczeń przez ponad 2000 lat.

C4.1. Trzy generacje komór oscylacyjnych

Analiza zasady działania komory oscylacyjnej ujawnia, że zrealizowanie tego urządzenia wcale nie wymaga aby jego kształt był dokładnie sześcienny. Przykładowo, zasada działania komór pierwszej generacji może też być zrealizowana w równoległościennym, w którym jedynie przekrój obiegu iskieł jest kwadratowy. Ponieważ jednak komory sześciennie będą najbardziej typowymi dla pierwszej generacji tych urządzeń - patrz podrozdział C7.1.2, dla uproszczenia rozważań w niniejszej monografii omawiana jest jedynie ich zasada działania i warunki operacyjne. W podobny jednak sposób jak w sześciennym, zasada ta może również zostać zrealizowana w kilku innych kształtach. Stosunkowo podobna do komory sześciennym będzie komora równoległościenna o przekroju kwadratowym. W jakiś więc czas po zbudowaniu **komory sześciennym**, na naszej planecie opracowana też zostanie **komora o kształcie równoległościennym** z kwadratowym przekrojem poprzecznym. Równoległościenny taki posiadał będzie cztery identyczne ścianki boczne w kształcie prostokątów, oraz dwie identyczne ścianki czołowe w kształcie kwadratów - np. patrz komora (M) na rysunku C9. Ta równoległościenna komora nie będzie jednak już typowa (patrz podrozdział C7.1.2) stąd jej zastosowanie ograniczało się będzie jedynie do kilku przypadków gdy jej kształt odmienny od sześciennym będzie absolutnie niezbędny. Najlepszym przykładem użycia takiego równoległościennym jest komora główna (M) w konfiguracji krzyżowej pokazanej na rysunku C9. Z uwagi na fakt, iż budowa i technologia wytwarzania komór oscylacyjnych o czterech ścianach bocznych, tj. w kształcie sześciennym oraz równoległościennym o przekroju kwadratowym, rozpracowane zostaną technicznie na naszej planecie jako pierwsze, nazywali je będziemy **"komorami pierwszej generacji"**.

Wygląd wszystkich komór pierwszej generacji będzie podobny. Jak to już opisano w poprzednim podrozdziale, będą one sprawiały wrażenie przeźroczystych kryształów o przekroju kwadratowym w płaszczyźnie rotacji ich iskieł. W środku będą one wypełnione złotymi iskrami jakby zamrożonymi w swoim migotaniu, a także gęstym polem magnetycznym przypominającym czarny dym.

Komory oscylacyjne pierwszej generacji będą zdolne do wypełniania ogromnej liczby różnorodnych funkcji. Ich tylko zgrubnemu omówieniu poświęcony jest cały podrozdział C9. Aby dać tu jakieś pojęcie o rozpiętości tych funkcji, to przykładowo są one zdolne do gromadzenia w swoim środku nielimitowanych ilości energii. Stąd na naszej planecie

całkowicie wyeliminują one dzisiejsze słupy wysokiego napięcia, linie przesyłowe, oraz transformatory elektryczności. Jeśli zaś użyte zostaną w napędzie magnokraftu, komory te będą służyć do wytwarzania sił nośnych, napędowych i manewrowych. Będą też podnosiły na statek wybrane obiekty (tj. działały jako efektywnie sterowalny dźwиг magnetyczny lub urządzenie zdalnego oddziaływania - patrz też podrozdział C7.3). Będą akumulowały zapasy energii statku (czyli stanowiły jakby jego "zbiorniki paliwa"). Wytwarzały promienie świetlne oświetlające wybrane obszary pod statkiem (czyli działały jak ogromne reflektory - patrz też podrozdział F1.3). Utrzymywały temperatury pomieszczeń statku na wymaganym i stałym poziomie (tj. działały jako klimatyzatory powietrza - patrz podrozdziały H6.1.3, C6.3 i F1.4). W bardziej zaś zaawansowanych wersjach tej komory, będą one zapewniały łączność telepatyczną (podrozdział F1.5). Ponadto będą także służyły ogromnej liczbie innych funkcji których wyjaśnienie wymagałoby znacznie dłuższych opisów.

Niestety, na określonym etapie rozwoju naszej cywilizacji, komory pierwszej generacji przestaną wypełniać wszystkie nakładane na nie wymagania. Szczególnie dwa czynniki będą tu decydujące, tj. (1) konieczność efektywnego napełniania komór energią, oraz (2) podjęcie budowy wehikułów teleportacyjnych jakie nałożą nowe wymagania dotyczące niezwykle ścisłego sterowania "przebiegiem w czasie" pulsowań wytwarzanego przez nie pola. (Przez "przebieg w czasie" należy rozumieć matematyczną funkcję $F=f(t)$ wyrażającą zależność zmian strumienia magnetycznego pola F od czasu t - np. patrz rysunek C7.) Aby sprostać tym wymaganiom budowa nowej, drugiej generacji komór oscylacyjnych musi zostać podjęta.

Komorami drugiej generacji nazywali będziemy komory zdolne do wyzwalania w sobie efektu telekinetycznego. Efekt ten nada im dwa atrybuty poprzednio nie występujące w komorach pierwszej generacji, tj. (1) będą one zdolne do samowzbudzenia swoich oscylacji pracując jako efektywne baterie telekinetyczne napełniające same siebie wymaganą ilością energii - po szczegóły patrz opisy z podrozdziałów K2.4 i L1, oraz (2) będą one zdolne do nadawania wyposażonemu w nie magnokraftowi zdolność to lotów w konwencji telekinetycznej - po szczegóły patrz podrozdziały B1, L1, i M6. Pierwszy z powyższych atrybutów wprowadzony może zostać już do komór oscylacyjnych o przekroju kwadratowym - patrz etap 10 procedury rozwojowej opisanej w podrozdziale C8.2. Jednakże łatwo przewidzieć, iż wytworzenie ciągu teleportacyjnego wprowadzi zaostrome wymagania sterownicze, jakich sprostanie wymusi aby zasada działania komór drugiej generacji realizowana była w komorze o kształcie **równoległościanu ośmiobocznego**. Komora taka posiadała więc będzie osiem identycznych ścianek bocznych w kształcie prostokątów, oraz dwie identyczne ścianki czołowe w kształcie ośmioboków równoramiennych. Niestety sterowanie tej komory oraz problemy techniczne związane z jej budową będą wielokrotnie bardziej złożone od sterowania i budowy komór o przekroju kwadratowym. Stąd też jej opracowanie będzie mogło nastąpić dopiero na znacznie wyższym etapie naszego rozwoju, na długo po opanowaniu technologii budowy i sterowania zwykłych komór o przekroju kwadratowym. Jednakże taka komora ośmioboczna dostarczać będzie pola magnetycznego jakiego charakterystyka znacznie przekroczy precyzję pola wytwarzanego przez komorę czworoboczną (sześcienną). Dla przykładu, pole stałe produkowane przez kapsułę dwukomorową złożoną z takich właśnie komór ośmiobocznych będzie wielokrotnie "bardziej stałe" niż pole stałe otrzymane ze zwykłej kapsuły dwukomorowej zawierającej sześcienną komorę (rozważ wpływ zwiększonej ilości członów w matematycznym ciągu Fouriera na wartość wynikową takiego ciągu).

Niezależnie od formowania efektu telekinetycznego, wysyłania telekinetycznego promienia podnoszącego (patrz opis z podrozdziału H6.2.1), oraz działania jako baterie telekinetyczne samozapelniające je wymaganą energią elektryczną, komory oscylacyjne drugiej generacji będą zdolne do wytwarzania jeszcze innego istotnego zjawiska. Będą one mianowicie pracowały również jako wysoko-efektywne nadajniki i odbiorniki telepatyczne, zdolne do zapewniania swoim użytkownikom natychmiastowej łączności telepatycznej z najodleglejszymi zakątkami wszechświata. Zasada ich działania przy tej dodatkowej funkcji

zrozumiana może zostać po przeanalizowaniu treści podrozdziałów H7.1, N2 niniejszej monografii.

Z wyglądu komory oscylacyjne drugiej generacji będą nieco podobne do komór oscylacyjnych pierwszej generacji, tyle tylko że ich geometria będzie nieco odmienna. Będą one miały kształt przeźroczystego kryształu o kształcie równoległościanu ośmiobocznego (czyli "dziesięciościanu"), zamiast - jak w przypadku komór oscylacyjnych pierwszej generacji - kryształu czworobocznego (w kształcie kostki sześciennego). Ich wygląd pokazano w części (b) rysunku C3. Z uwagi na warunki konstrukcyjne i użytkowe opisane w podrozdziale C7.1.2, proporcje wymiarowe D/H tych równoległościanów (tj. stosunek średnicy D okręgu opisanego na ich czołach do wysokości H całej komory) w typowych komorach drugiej generacji będą ściśle określone i równe $D/H=1$ (patrz rysunek C8). Podobnie jak komory pierwszej generacji, komory oscylacyjne drugiej generacji będą również wypełnione iskrami rotującymi dookoła osi magnetycznej "m" komory (tj. dookoła obwodu jej ośmiobocznych ścian czołowych - patrz część (b) rysunku C3).

Po komorach drugiej generacji kolej przyjdzie na budowę komór trzeciej generacji. Ich podstawowym atrybutem będzie iż potrafią one spowodować zmiany w upływie czasu (patrz magnetyczna interpretacja czasu opisana w podrozdziałach H9.1 i M1). Już obecnie daje się przewidzieć, iż **komory trzeciej generacji** będą oparte na komorach szesnastobocznych. Z wyglądu będą więc one przypominać rodzaj niemal okrągłego pręta, jakiego wnętrze widoczne poprzez przeźroczyste ścianki będzie przypominało wnętrze komór poprzednich generacji (tj. także wypełnione będzie rotującymi iskrami, tyle że mniejszymi i bardziej równomiernie rozproszonymi w objętości tych komór). Ich budowa zostanie zainicjowana z chwilą gdy nasza cywilizacja rozpocznie prace nad wehikułami czasu.

Komory oscylacyjne trzeciej generacji będą zdolne do wytwarzania pełnego zakresu zjawisk omawianych w niniejszej monografii. Niezależnie od zdolności do powodowania zmian w upływie czasu będą one też zdolne do formowania efektu telekinetycznego, do działania jako urządzenia zdalnego telekinetycznego "beaming up", do pracy jako baterie telekinetyczne, a także do działania jako telepatyczne stacje nadawczo-odbiorcze. Na dodatek do tego będą one zdolne do formowania wszystkich efektów powodowanych przez komory oscylacyjne pierwszej generacji (np. siły odpychnia magnetycznego, oświetlenia, termicznej klimatyzacji pomieszczeń, itp.).

Z wyglądu komory oscylacyjne trzeciej generacji również będą nieco podobne do komór oscylacyjnych pierwszej i drugiej generacji. Tyle tylko, że ich geometria będzie nieco odmienna. Mianowicie będą one przyjmowały formę przeźroczystego kryształu o kształcie równoległościanu szesnastobocznego (czyli "osiemnastościanu"), zamiast - jak w przypadku komór oscylacyjnych pierwszej i drugiej generacji - kryształu sześciennego lub dziesięciościennego. Przy niezbyt dokładnym obejrzeniu będą one sprawiały wrażenie niemal okrągłego cylindra którego średnica jest równa wysokości, tj. $D=H$ (patrz rysunki C3, C8 i C11). Ponownie z uwagi na warunki konstrukcyjne i użytkowe opisane w podrozdziale C7.1.2, proporcje wymiarowe D/H tych równoległościanów (tj. stosunek średnicy D okręgu opisanego na ich czołach do wysokości H całej komory) w typowych komorach będą ściśle określone i równe $D/H=1$ (patrz część (3s) na rysunku C8). Podobnie jak komory pierwszej i drugiej generacji, komory oscylacyjne trzeciej generacji będą również wypełnione iskrami rotującymi dookoła osi magnetycznej komory (tj. dookoła obwodu jej szesnastobocznych ścian czołowych). Tyle tylko że iskry rotujące wzdłuż obwodów komór trzeciej generacji będą jeszcze bardziej jednorodne, delikatne i równomierniej rozłożone jak iskry w komorach drugiej (i pierwszej) generacji.

* * *

Jak to z powyższego można wywnioskować, kształt zewnętrzny jaki przyjmie dana komora oscylacyjna będzie bezpośrednim wskaźnikiem poziomu zaawansowania technicznego cywilizacji która urządzenie to używa. Stąd też jest istotnym aby znać owe kształty bowiem umożliwi to nam rozpoznawanie poziomu rozwojowego do którego dana cywilizacja należy, a także zasadę lotu używanych przez nią wehikułów magnokrafto-

podobnych (tj. czy są to wehikuly magnetyczne, teleportacyjne, czy też wehikuly czasu - patrz podrozdziały M6, oraz T1 do T4).

C5. Matematyczny model komory oscylacyjnej

Nasza obecna znajomość zjawisk elektrycznych i magnetycznych umożliwia nam wyprowadzenie równań wyrażających związki pomiędzy wymaganymi wartościami oporności, indukcyjności i pojemności komory oscylacyjnej w kształcie sześcienu. Następne złożenie tych równań z sobą i ich analiza umożliwi wnioskowanie o zachowaniu się tego urządzenia. Dla uproszczenia rozważań wszystkie analizy wykonane zostaną wyłącznie dla komór o kształcie kostek sześciennych, stąd interpretacja uzyskanych wyników dla komór o innych kształtach pozostawiona zostanie do uznania czytelników.

Niniejszy podrozdział opisuje komorę oscylacyjną w języku matematyki. Dla przyszłych badaczy tego urządzenia dostarcza on więc istotnych podstaw interpretacyjnych. Jednakże dla czytelników mniej zorientowanych matematycznie, może on popsuć przyjemność zapoznawania się z tą monografią. Dlatego też tym z czytelników, u których wzory matematyczne wywołują nawrót senności, rekomendowałbym przejście z tego miejsca bezpośrednio do czytania podrozdziału C6.

C5.1. Oporność komory oscylacyjnej

Ogólna postać równania na oporność "R" dowolnego opornika o przekroju poprzecznym "A" i długości "l" jest jak następuje:

$$R = l \cdot (\Omega / A)$$

W równaniu tym "Ω" reprezentuje oporność właściwą materiału z jakiego dany opornik jest wykonany. W naszym przypadku będzie to maksymalna oporność gazu dielektrycznego jaki wypełnia komorę oscylacyjną, wyznaczona dla początkowej chwili wyładowania elektrycznego. Z kolei **operatory "*" oraz "/" zapożyczone z programowania komputerów oznaczają "mnożenie" oraz "dzielenie"**.

Jeśli w powyższym równaniu zastąpić jego zmienne przez poszczególne wartości wyznaczone dla komory oscylacyjnej, tj. $l = a$ oraz $A = a$ (porównaj z rysunkiem C1 "b"), wtedy otrzymamy:

$$R = \Omega / a \tag{C1}$$

Otrzymane w ten sposób równanie opisuje oporność elektryczną "R" sześcienniej komory oscylacyjnej, jaka jest funkcją wymiarów "a" jej ścianki bocznej.

C5.2. Indukcyjność komory oscylacyjnej

Dokładne wyznaczenie indukcyjności komory oscylacyjnej jest niezwykle trudnym i kompleksowym zadaniem. Jego skompletowanie z pełną dokładnością przekracza moje możliwości. Także wszyscy eksperci w tym zakresie jakich konsultowałem nie potrafili pomóc (być może któryś z czytelników znajdzie sposób jak rozwiązać ten problem - w takim przypadku chętnie zapoznałbym się z tokiem samego wyprowadzenia i końcowym wynikiem). Nie będąc w stanie znaleźć dokładnego rozwiązania problemu, zdecydowałem się wprowadzić założenie upraszczające. Aby uzasadnić to założenie warto wspomnieć, iż równanie (C2) na indukcyjność komory wyprowadzone w taki uproszczony sposób użyte

będzie w dalszej części monografii tylko jednokrotnie, kiedy znaczenie współczynnika "s" (patrz równanie (C5)) jest interpretowane. Stąd przyjęte tu uproszczenie nie wpływa na żadne z istotnych równań niniejszej monografii.

W uproszczonych wyprowadzeniach indukcyjności komory **przyjęto założenie**, iż jednostkowa induktancja pęku iskier (tj. induktancja odniesiona do jednostki nominalnej długości iskry) będzie równa induktancji takiego samego odcinka cewki. Owo założenie umożliwia więc wykorzystanie szeroko znanego równania na induktancję "L" selenoidu (patrz książka [1C5.2] pióra David Halliday et al, "Fundamentals of Physics", John Wiley & Sons, 1966):

$$L = \mu \cdot n^2 \cdot l \cdot A$$

Kiedy w równaniu tym dokonamy podstawień: $n = p/a$, $l = a$ i $A = a$ (gdzie "p" jest liczbą segmentów wydzielonych w każdej elektrodzie, podczas gdy "a" jest długością boku każdej ze ścian komory sześciennnej), wtedy uzyskane zostanie uproszczone równanie na induktancję "L" komory oscylacyjnej:

$$L = \mu \cdot p \cdot a \tag{C2}$$

W celu usprawiedliwienia podjętego tu uproszczenia można wykazać teoretycznie, iż jednostkowa inercja elektryczna (induktancja) pęku iskier będzie znacznie większa od takiej inercji w odpowiadającym odcinku przewodnika. Poświadczenia tego faktu dostarcza analiza mechanizmu zjawiska inercji. Inercja bowiem ujawnia swoje działanie gdy dany ruch obejmuje odwracalne zjawiska, lub transformowalne substancje, jakie w początkowym stadium rozwoju danego ruchu absorbują energię, aby potem ją wyzwolić w stadium zanikania tego samego ruchu. Im większa liczba takich odwracalnych zjawisk i substancji objęta zostaje danym ruchem, oraz im większa jest pochłaniana przez nie energia, tym wynikowa inercja jest większa. Pęk iskier przeskakujących w gazie, w każdym aspekcie wykazuje większy potencjał dla powodowania inercji niż prąd przepływający przez przewodnik. Pierwszą przyczyną dla tego stanu rzeczy jest bardziej efektywna absorpcja oraz uwalnianie energii przez iskry, następujące ponieważ:

(a) Szybkość elektronów w iskrze może być większa niż szybkość elektronów w przewodniku.

(b) Poszczególne iskry danego pęku mogą przeskakiwać w bliższych odległościach od siebie niż zwoje przewodnika w cewce, ponieważ nie będą one wymagały warstewek izolacyjnych do oddzielania ich od siebie.

Druga przyczyna dla tej wysokiej inercji iskier w gazie wynika z obejmowania przez nie większej liczby odwracalnych zjawisk - jakie nie występują wcale podczas przepływu prądu przez zwoje przewodnika. Zjawiska te to:

(c) Jonizacja otaczającego gazu. Jonizacja ta, dzięki późniejszemu oddawaniu zaabsorbowanej początkowo energii, dodatkowo zwiększa inercję w momencie zaniku iskier.

(d) Powodowanie ruchu ciężkich jonów, jakich masa absorbuje i potem uwalnia energię kinetyczną (znacznie większą od energii lekkich elektronów poruszających się w metalowym przewodniku).

(e) Zainicjowanie zjawisk hydrodynamicznych (np. ciśnienia dynamicznego gazu) jakie także będą powodowały przemieszczanie się ładunków elektrycznych oraz zwrot energii w momencie zaniku iskier.

Powyższe przesłanki teoretyczne nie powinny być trudne do praktycznego zweryfikowania za pomocą eksperymentów opisanych w podrozdziale C8.2 (np. etap 1c).

C5.3. Pojemność komory oscylacyjnej

Jeśli użyjemy dobrze znanego równania na pojemność "C" kondensatora płytkowego (o dwóch równoległych elektrodach o powierzchni "A" i wzajemnej odległości "l"), o następującej postaci:

$$C = \epsilon \cdot A / l$$

i następnie podstawimy do niego wartości: "A = a²" i "l = a", da to nam końcowe równanie na pojemność "C" komory oscylacyjnej:

$$C = \epsilon \cdot a \quad (C3)$$

(tj. pojemność "C" komory oscylacyjnej jest równa wartości stałej dielektrycznej "ε" dla gazu wypełniającego tą komorę, pomnożonej przez długość boku "a" tej komory).

C5.4. Współczynnik motoryczny iskier i jego interpretacja

Każda z zależności (C1), (C2) i (C3) opisuje tylko jeden wybrany parametr komory oscylacyjnej. Z drugiej jednak strony, byłoby wysoce użytecznym posiadać pojedynczy złożony współczynnik jaki byłby w stanie wyrazić jednocześnie wszystkie elektromagnetyczne i konstrukcyjne właściwości danej komory oscylacyjnej. Wprowadźmy teraz taki współczynnik, nazywając go "współczynnikiem motorycznym iskier". Jego równanie definicyjne jest jak następuje:

$$s = p \cdot \frac{R}{2} \cdot \sqrt{\frac{C}{L}} \quad (C4)$$

Zauważ, że po jego zapisaniu w notacji komputerowej, w której symbol "*" oznacza mnożenie, symbol "/" oznacza dzielenie, zaś symbol "sqrt()" oznacza pierwiastek kwadratowy z argumentu podanego w nawiasie (), owo równanie (C4) przyjmuje następującą postać: $s = p \cdot (R/2) \cdot \text{sqrt}(C/L)$.

Proszę zwrócić uwagę, iż zgodnie z powyższym równaniem definiującym, współczynnik "s" jest bezwymiarowy.

Niezależnie od powyższego równania definiującego, współczynnik "s" posiada również opis interpretacyjny. Opis ten może zostać uzyskany gdy w równaniu (C4) zmienne R, L, i C zostaną zastąpione przez wartości wyrażone równaniami (C1), (C2) i (C3). Kiedy to zostanie dokonane, wtedy otrzymane zostanie następujące interpretacyjne równanie na "s":

$$s = \frac{1}{2a} \cdot \Omega \cdot \sqrt{\frac{\epsilon}{\mu}} \quad (C5)$$

Zauważ, że po jego zapisaniu w notacji komputerowej, w której symbol "*" oznacza mnożenie, symbol "/" oznacza dzielenie, zaś symbol "sqrt()" oznacza pierwiastek kwadratowy z argumentu podanego w nawiasie (), owo równanie (C5) przyjmuje następującą postać: $s = (1/(2 \cdot a)) \cdot \Omega \cdot \text{sqrt}(\epsilon/\mu)$.

Równanie (C5) ujawnia, że współczynnik "s" doskonale reprezentuje aktualny stan wszystkich warunków otoczeniowych w jakich zachodzą wyładowania iskrowe w komorze, a jakie wyznaczają ich przebieg i efektywność. Tak więc opisuje ono rodzaj i konsystencję gazu

użytego w komorze jako dielektryk, oraz aktualne parametry pod jakimi gaz ten się znajduje. Również opisuje ono wymiary komory. Stąd współczynnik "s" stanowi doskonały parametr zdolny do dokładnego opisu aktualnej sytuacji roboczej panującej w komorze w danym momencie czasu.

Wartość współczynnika "s" może być sterowana zarówno na etapie konstrukcji komory, jak i na etapie jej eksploatacji. Na etapie konstrukcji jest to uzyskiwane poprzez zmiany we wymiarze boku "a" komory sześcienną. Natomiast na etapie eksploatacji wymaga to zmian w ciśnieniu gazu dielektrycznego zawartego w komorze, lub zmienienia jego kompozycji. W obu przypadkach taka zmiana ciśnienia lub kompozycji tego gazu wpłynie na wartość stałych Ω , μ i ϵ , opisujących jego własności elektryczne. (Odnótuj że stałe " Ω ", " μ ", oraz " ϵ ", posiadają następujące interpretacje: Ω = oporność elektryczna gazu dielektrycznego w komorze wyznaczona w momencie początku przeskoku iskry w [Ohm*metr], μ = przenikalność magnetyczna gazu dielektrycznego w [Henry/metr], ϵ = stała dielektryczna dla gazu wypełniającego komorę w [Farad/metr].)

C5.5. Warunek zaistnienia oscylacji we wnętrzu komory

Z elektrycznego punktu widzenia komora oscylacyjna reprezentuje typowy obwód RLC. Badania dokonane na sieciach elektrycznych (electric networks) wyznaczyły dla takich obwodów warunek jakiego spełnienie jest wymagane aby dany obwód po jednorazowym naładowaniu go elektrycznością, dostarczył oscylacyjnej odpowiedzi (tj. zareagował poprzez wytworzenie ciągu oscylacji). Warunek ten, matematycznie przedstawiony w książce [1C5.5] pióra Hugh H. Skilling, "Electric Network" (John Willey & Sons, 1974), przyjmuje następującą postać:

$$R < 4 \cdot L/C$$

Jeśli powyższą relację przekształcić i następnie jej zmienne zastąpić równaniem (C4), wtedy przyjmie ona następującą formę końcową:

$$p > s \tag{C6}$$

Powyższy warunek opisuje więc konstrukcyjne wymaganie na liczbę segmentów "p" wydzielonych w elektrodach komory oscylacyjnej, w odniesieniu do warunków otoczenia "s" panujących w obszarze roboczym tej komory poprzez który dane iskry muszą przeskakiwać. Jeśli ten warunek zostanie wypełniony, wtedy iskry produkowane w komorze oscylacyjnej będą posiadały oscylacyjny charakter.

Aby zinterpretować warunek (C6), możliwy zakres wartości przyjmowanych przez współczynnik "s" powinien zostać rozpatrzony (porównaj warunek (C6) z równaniem (C5)).

C5.6. Okres pulsowań pola komory oscylacyjnej

Z obwodów typu "RLC" wiadomo, iż okres "T" ich oscylacji wyraża następujące równanie:

$$T = \frac{2 \cdot \Pi}{\sqrt{\frac{1}{L \cdot C} - \left(\frac{R}{2 \cdot L}\right)^2}} = 2 \cdot \Pi \sqrt{\frac{L \cdot C}{1 - \frac{R^2 C}{4 \cdot L}}}$$

Zauważ, że po jego zapisaniu w notacji komputerowej, w której symbol "-" oznacza odejmowanie, symbol "*" oznacza mnożenie, symbol "/" oznacza dzielenie, symbol "R**2" oznacza "R" podniesione do potęgi "2", zaś symbol "sqrt()" oznacza pierwiastek kwadratowy z argumentu podanego w nawiasie (), powyższe równanie przyjmuje następującą postać: $T = (2 \cdot \Pi) / (\text{sqrt}(1/(L \cdot C) - (R/(2 \cdot L))^2)) = 2 \cdot \text{sqrt}(L \cdot C / (1 - ((R^2 \cdot C)/(4 \cdot L)))$.

Jeśli wielkości definiujące współczynnik "s" z równania (C4) w powyższym równaniu zastąpią kombinację parametrów R, L, i C, podczas gdy równania (C1) i (C3) dostarczą wartości dla R i C, wtedy ów okres pulsacji zostaje opisany następującym równaniem:

$$T = \frac{\Pi \cdot \frac{p}{s} \cdot \Omega \cdot \epsilon}{\sqrt{1 - \left(\frac{s}{p}\right)^2}} \quad (C7)$$

Zauważ, że po jego zapisaniu w notacji komputerowej, w której symbol "-" oznacza odejmowanie, symbol "*" oznacza mnożenie, symbol "/" oznacza dzielenie, symbol "**" oznacza podnoszenie do podanej potęgi, zaś symbol "sqrt()" oznacza pierwiastek kwadratowy z argumentu podanego w nawiasie (), owo równanie (C7) przyjmuje następującą postać:

$$T = (\Pi \cdot (p/s) \cdot \Omega \cdot \epsilon) / \text{sqrt}(1 - (s/p)^2).$$

Końcowe równanie (C7) nie tylko że wyraża od czego w komorze oscylacyjnej zależy jest okres jej pulsowań "T", ale także pokazuje praktycznie w jaki sposób wartość tego okresu "T" może być sterowana. Będzie ono więc wysoce użyteczne dla zrozumienia amplifikującej zasady sterowania komorą opisaną w podrozdziale C6.5.

Znając okres pulsowań "T" pola magnetycznego komory, możliwe jest też łatwe wyznaczenie częstości pulsowań "f" tego pola. Szeroko bowiem znana współzależność pomiędzy tymi wielkościami jest jak następuje:

$$f = 1/T \quad (C8)$$

Oczywiście zgodnie z powyższym równaniem (C8), sterowanie częstością "f" pulsowań pola w komorze będzie odbywać się poprzez sterowanie okresem "T" tego pola uzyskiwane poprzez wykorzystanie interpretacji wzoru (C7).

C6. Jak komora oscylacyjna eliminuje wady elektromagnesów

Działanie komory oscylacyjnej zostało ukształtowane w taki sposób, że wszystkie wady wrodzone elektromagnesów zostają w niej całkowicie wyeliminowane. Opisy jakie nastąpią w dalszych częściach niniejszego podrozdziału zaprezentują istotę wyeliminowania pięciu najważniejszych wad elektromagnesów, wymienionych i omówionych w punktach #1 do #5 podrozdziału C1.

C6.1. Neutralizacja sił elektromagnetycznych

Jedną z najistotniejszych wad elektromagnesów była siła odchylająca powstająca w ich zwojach (opisana w punkcie #1 podrozdziału C1). Siła ta w efekcie końcowym prowadzi do eksplozji tych urządzeń przy przekroczeniu przez nie określonego wydatku granicznego. W komorze oscylacyjnej ta sama siła zostaje całkowicie zneutralizowana. Unikalna bowiem zasada działania komory powoduje powstanie w niej nie jednej, a aż dwóch nawzajem przeciwnych sił, tj. (1) Coulomb'owskiego przyciągania się przeciwnych ścianek, oraz (2) owej elektromagnetycznej siły odchylającej (tj. tej samej jaka rozrywała elektromagnesy). Obie te siły, działając jedna na drugą, nawzajem się więc zneutralizują. Niniejszy podrozdział prezentuje zasadę na jakiej neutralizacja ta nastąpi.

Siły Coulomb'owskie powstają w efekcie wzajemnego przyciągania się obu ładunków elektrostatycznych "+q" i "-q" o równych wartościach ale przeciwnych znakach, zgromadzonych na obu nawzajem przeciwnych sobie ściankach komory. Siły te ściskają komorę dośrodkowo, starając się ją zgnieść. Z kolei siły odchylające wytwarzane są wskutek oddziaływania przeskakujących iskier z polem magnetycznym panującym w komorze. Siły te starają się rozzerwać komorę odśrodkowo. Stąd też możliwym jest takie dobranie konstrukcji i warunków operacyjnych komory aby oba te układy sił nawzajem się znosiły (zerowały). Po takim dobraniu ścianki komory będą z taką samą siłą ściskane dośrodkowo przez siły Coulomb'owskie, jak rozrywane odśrodkowo przez siły odchylające. Ponieważ oba układy sił działają również wzajemnie na siebie, ich wynikowy efekt będzie równy zero, czyli taki sam jaki byłby w przypadku braku jakichkolwiek sił działających w danej komorze. W rezultacie końcowym więc, struktura fizyczna komory będzie uwolniona od konieczności przeciwstawiania się jakimkolwiek siłom elektromagnetycznym.

Rysunek C4 pokazuje mechanizm wzajemnej neutralizacji obu tych układów sił. Dla uproszczenia, wszystkie przebiegi zjawisk zachodzących w komorze pokazane zostały tam jako zjawiska liniowe, niezależnie od tego jak zachodzą one w rzeczywistości. Daje się jednakże wydedukować, że w rzeczywistości zjawiska te muszą być symetryczne. Oznacza to iż jeśli, dla przykładu, prąd w iskrach zmieni się w określony sposób, wtedy również i potencjał na elektrodach musi się zmienić dokładnie w taki sam sposób. Stąd "zmiany w czasie" sił analizowanych tutaj wykazują swoisty wrodzony mechanizm samo-regulacyjny. W mechanizmie tym przebieg (ale nie ilość) jednego zjawiska zawsze podąża za przebiegiem drugiego ze zjawisk. W ten sposób niezależnie jakie są rzeczywiste zmiany w czasie dla omawianych tu zjawisk siłowych, zasada ich wzajemnej neutralizacji omówiona na przykładzie przebiegu liniowego będzie także ważna dla ich rzeczywistych przebiegów.

Część (a) rysunku C4 pokazuje cztery podstawowe fazy formowania pełnego cyklu działania komory. Opis tych faz przedstawiony był już w podrozdziale C3.3 tego rozdziału. Istotnym dla każdej z faz jest, iż jednocześnie współistnieją w niej dwa pęki iskier, pierwszy z których, na rysunku C4 (a) pokazany linią ciągłą, przekazuje energię z pola elektrycznego do pola magnetycznego (są to więc iskry aktywne). Natomiast drugi z pęków (na rysunku C4 (a) pokazany za pomocą linii przerywanej) w tym samym czasie konsumuje pole magnetyczne i wytwarza pole elektryczne (iskry inercyjne).

Część (b) rysunku C4 pokazuje odpowiadające tym iskrom zmiany w ładunkach "q" na prawej R (tj. right), lewej L (tj. left), przedniej F (tj. front) i tylnej B (tj. back) elektrodzie, następujące w każdej z czterech faz działania komory. Owe ładunki wytwarzają siły Coulomb'owskie jakie przyciągają dośrodkowo położone przeciwległe elektrody. W tej części rysunku uwidoczniło się, podczas gdy jedna para elektrod osiąga maksimum swej różnicy potencjałów - inicjując wyładowanie pomiędzy nimi, równocześnie druga z par jest w równowadze swoich potencjałów. Następnie równocześnie ze wzrostem prądu wyładowania przepływającego pomiędzy pierwszą parą elektrod, przeciwstawne ładunki elektrostatyczne zgromadzone na drugiej parze elektrod również wzrastają. Wiadomo iż siły odchylające jakie rozrywają komorę odśrodkowo rosną wraz ze wzrostem wartości prądu wyładowania. Siły te dla przeskoku iskry pomiędzy elektrodami jednej pary powodują napór iskier na elektrody

drugiej pary. Z drugiej strony siły Coulomb'owskie wzajemnego przyciągania się tej drugiej pary naprzemianległych elektrod także wzrastają. Dzięki temu mechanizmowi oba przeciwstawne sobie rodzaje sił rosną w tym samym tempie.

Część (c) rysunku C4 pokazuje zmiany w elektromagnetycznych siłach odchyłających $M=i$ a B , starających się wypchnąć poszczególne pęki iskier z zasięgu pola komory. Ponieważ siły te są proporcjonalne do produktu prądu iskier "i" oraz gęstości pola magnetycznego " $B=F/(a^2)$ ", maksymalna wielkość wywołanego nimi rozrywania komory wypadnie w momencie czasowym kiedy wyładowujące daną iskrę elektrody osiągną równowagę swoich potencjałów. Jednakże właśnie w tym samym momencie czasu druga para elektrod, do których owe iskry są dociskane, osiąga maksimum swojej różnicy potencjałów (porównaj z częścią (b) tego rysunku) a co za tym idzie i maksimum swoich sił Coulomb'owskiego przyciągania. W swoich maksimach oba rodzaje sił także więc nawzajem się kompensują.

W części (d) rysunku C4 pokazano mechanizm wzajemnej neutralizacji sił opisanych uprzednio. Górna połowa wykresu z tej części rysunku pokazuje zmiany w siłach odchyłających "T", jakie starają się rozerwać komorę. Siły te wywoływane są przez wzajemne oddziaływanie pola komory i prądu iskier (porównaj z częścią (c) tego rysunku). Dolna połowa wykresu z tej części (d) rysunku C4 pokazuje zmiany w siłach ściskających "C". Siły te są wywoływane przez Coulomb'owskie przyciąganie pomiędzy naprzemianległymi elektrodami jakie akumulują przeciwne ładunki elektrostatyczne "q" (porównaj z częścią (b) rysunku C4). Zauważ, że kiedykolwiek w komorze pojawia się siła rozrywająca "T" (np. z pędu iskier S_{B-F}), zawsze równocześnie formowana jest przeciwdziałająca jej siła ściskająca "C" (np. z Coulomb'owskiego przyciągania ładunków q_{R-L}). Obie te siły działają w przeciwstawnych kierunkach oraz zmieniają się według tych samych przebiegów w czasie. Stąd też obie one neutralizują się nawzajem.

Oczywiście jest zrozumiałe, iż opisana tu neutralizacja sił, od początku wykazująca symetryczność przebiegów siłowych (jak to wykazano już poprzednio), ciągle wymaga dopasowania swych wartości. Stąd też konieczne będą eksperymenty naukowe podczas budowy komory oscylacyjnej jakie pozwolą na dobranie takich parametrów konstrukcyjnych i eksploatacyjnych tego urządzenia jakie spowodują kompletne zrównoważenie się obu tych przeciwstawnych sił. Wynikiem tych badań będzie skompletowanie komory, w której wytwarzanie pola magnetycznego nie będzie ograniczane działaniem żadnego z omawianych tu rodzajów sił. Pole to będzie więc mogło wzrastać do teoretycznie Nielimitowanych niczym wartości, wielokrotnie przekraczając nawet "strumień startu".

C6.2. Niezależność wytwarzanego pola od ciągłości i efektywności dostawy energii

Jednym z najbardziej podstawowych atrybutów każdego układu oscylującego jest zdolność do absorbowania energii dostarczanej do niego w sposób nieciągły. Przykładem takiej nieciągłej dostawy jest dziecko na huśtawce. Huśtawki tej nie musimy wszakże ciągle popychać. Wystarczy iż dodamy jej energii raz na jakiś czas, a mimo to będzie ona kontynuowała swój ruch oscylacyjny w sposób ciągły. Powyższe praktycznie oznacza, że energia raz dostarczona do komory oscylacyjnej zostanie uwięziona w niej na tak długo aż zaistnieją zewnętrzne okoliczności jakie spowodują jej wycofanie. Jak to zaś zostanie wyjaśnione w podrozdziale C6.3.1 takie okoliczności zaistnieją tylko jeśli komora zostanie użyta do wykonywania jakiejś zewnętrznej pracy.

Innym istotnym atrybutem układów oscylujących jest ich zdolność superpozycji czyli możliwość zmiany poziomu zawartej w nich energii na drodze okresowego dodawania dalszych porcji energii do zasobów już w nich zgromadzonych. W poprzednim przykładzie huśtawki, aby spowodować wyniesienie dziecka na określoną wysokość wcale nie jest koniecznym nadanie naraz huśtawce całej wymaganej przez nią energii. Wystarczy bowiem popychać ją po troszeczkę przez dłuższy okres czasu, dodając energii stopniowo. Następstwem tego atrybutu jest, że komora oscylacyjna nie będzie wymagała dostarczenia jej

całego zasobu energii w jednym impulsie. Stąd dostawa energii do tego urządzenia może być stopniowa i rozłożona na dłuższy okres czasu.

Oba omawiane atrybuty razem dostarczają nam praktycznej drogi dla dostarczania do komory każdej ilości energii jaka może być wymagana przez produkowane przez nią pole magnetyczne, bez wprowadzania żadnych wymagań czy ograniczeń odnośnie źródła lub linii przesyłowej jakie użyte zostają w celu tego dostarczania.

Aby dopomóc nam w uświadomieniu sobie przewagi powyższego sposobu dostarczania energii do komory oscylacyjnej nad sposobem wymaganym dla elektromagnesów, użyjmy następującego przykładu. Dziecko na huśtawce i potężny atleta oboje starają się wydzwignąć spory ciężar na określoną wysokość. Dziecko czyni to niemalże bez wysiłku poprzez akumulowanie energii wychyłu podczas kolejnych oscylacji. Natomiast atleta musi użyć całej swojej mocy i ciągle cel może okazać się dla niego nieosiągalny.

C6.3. Eliminacja strat energii

Iskry są dobrze znane ze swojej wrodzonej zdolności do rozpraszania energii. Nie ma więc wątpliwości, iż taka intensywna cyrkulacja iskier, jak ta występująca w komorze oscylacyjnej, musi zamieniać dużą ilość elektryczności na ciepło. W zwyczajnym urządzeniu taka zamiana byłaby powodem znacznych strat energii. Jednakże podczas działania komory oscylacyjnej wystąpią unikalne warunki jakie umożliwią powrotną transformację energii cieplnej w elektryczność. Owa transformacja pozwoli na odzyskanie i ponowne zakumulowanie w formie przeciwnych ładunków elektrostatycznych całej energii poprzednio rozproszonej w postaci ciepła wydzielanego przez iskry. Tak więc w komorze oscylacyjnej współistnieć będą dwa równoczesne procesy: (1) rozpraszanie energii poprzez zamienianie części energii elektrycznej iskier na ciepło, oraz (2) odzyskiwanie energii poprzez bezpośrednią zamianę energii cieplnej w pole elektrostatyczne. Oba te procesy będą nawzajem neutralizowały efekty swego działania. Stąd bez względu na to ile wyniesie rozpraszanie energii przez poszczególne iskry, komora oscylacyjna jako całość zupełnie wyeliminuje ich straty energetyczne. Jako więc wynik końcowy takiej eliminacji, cała energia dostarczona do tego urządzenia będzie w nim zachowana na zawsze, jeśli oczywiście nie zostanie ona zużyta na wykonywanie przez komorę jakiejś pracy zewnętrznej.

W komorze oscylacyjnej współistnieją aż trzy elementy jakie w takiej samej konfiguracji i wartościach nie występowały jeszcze w żadnym z poprzednio budowanych przez nas urządzeń. Są to: silne pulsujące pole magnetyczne, elektrody, oraz gaz dielektryczny. Na dodatek w czasie działania tego urządzenia elementy te przyjmują stany jakie, zgodnie z moimi teoriami opisanymi w monografiach [1a], [3], [3/2], [6] i [6/2] wymagane będą do zaistnienia dotychczas jeszcze mało poznanego zjawiska, zwanego "efekt telekinetyczny" - patrz jego opis przytoczony w podrozdziale H6.1. Użycie efektu telekinetycznego do bezpośredniej zamiany ciepła w elektryczność uzależnione jest bowiem właśnie od współistnienia pola magnetycznego jakiego linie sił są przyspieszane i opóźniane, elektrod których ładunki fluktuują, oraz zjonizowanego gazu. Opis dosyć złożonej teorii stojącej za efektem telekinetycznym, sposobów jego technicznego wyzwania, oraz urządzeń energetycznych już zbudowanych na Ziemi jakie wykorzystują go w celach bezpośredniej zamiany energii cieplnej w elektryczność, wymaga długich wyjaśnień. Stąd też czytelnikom zainteresowanym w dokładniejszym poznaniu tego zjawiska rekomenduję zapoznanie się albo z podrozdziałem H6.1 i rozdziałami K i L niniejszej monografii, albo z następnym wydaniem monografii z serii [6] całkowicie poświęconej temu tematowi. W tym miejscu jednak pragnąłbym poinformować, że już potwierdzone eksperymentalnie działanie efektu telekinetycznego czyni z niego "zjawisko stanowiące odwrotność tarcia". Podobnie bowiem jak tarcie powoduje samoczynne konsumowanie ruchu oraz wytwarzanie ciepła, efekt telekinetyczny powoduje samoczynne konsumowanie ciepła oraz wytwarzanie ruchu. W przypadku więc komory oscylacyjnej będzie on użyty dla transformowania ciepła

generowanego przez iskry w ruch ładunków elektrycznych jakie spowodują wzrost potencjałów elektrostatycznych na jej elektrodach.

Zasadę całkowitego odzyskiwania energii cieplnej komory poprzez wykorzystywanie działania efektu telekinetycznego wyjaśniono w podrozdziale H6.1.3. Z uwagi jednak na utrzymywanie ciągłości rozważań podsumujemy tą zasadę krótko również i w tym miejscu. Efekt telekinetyczny umożliwia kontrolowane wyzwalać dwóch przeciwnych zjawisk termicznych powodujących m.in. naprzemienne emitowanie tzw. "światła pochłaniania" oraz "światła wydzielania". Podczas wyzwalać pierwszego z tych zjawisk energia cieplna otoczenia przetransformowana może zostać bezpośrednio w ruch, w drugim zaś zjawisku ruch przetransformowany może zostać bezpośrednio w energię cieplną. Kierunek i intensywność tych transformacji "ciepło/ruch" zależą od kierunku oraz od wartości wektora chwilowych przyspieszeń lub opóźnień linii sił pulsującego pola magnetycznego przenikającego daną objętość komory (a ściślej od wzajemnych relacji i zorientowania tego wektora względem chwilowych wektorów ruchu ładunków elektrycznych w komorze). Przy odpowiednim więc doborze przebiegu krzywych chwilowych pulsowań pola, a ściślej po właściwej desymetryzacji tych pulsowań, temperatura komory może być utrzymywana na stałym, niezmiennym i sterowalnym poziomie - patrz też opisy z podrozdziału K2.4. Zasada tego utrzymywania opiera się na takim sterowaniu przebiegiem krzywej chwilowych zmian pola magnetycznego wytwarzanego w danej komorze, aby poszczególne pół-pulsy tego pola wyzwalały efekt telekinetyczny powodujący odpowiednie przyspieszanie (lub opóźnianie) rotacji ładunków elektrycznych iskier kosztem energii termicznej zawartej w komorze. W ten sposób całe ciepło strat energetycznych zaistniałych w rezultacie przeskoku iskier, zamieniane będzie przez kontrolowany efekt telekinetyczny w ruch ładunków składających się na owe iskry. W rezultacie więc efekt telekinetyczny przetransformuje ciepło tracone przez iskry elektryczne na kontrolowany ruch ładunków elektrycznych, utrzymując w ten sposób temperaturę komory na stałym i z góry zadany poziomie.

Zdaję sobie sprawę, że moje stwierdzenia zawarte w niniejszym podrozdziale zapewne mogą wywołać opozycję ze strony osób dotychczas nieobznajomionych z efektem telekinetycznym. Stąd też dla użytku owych osób, w podrozdziale jaki nastąpi przytoczone zostaną argumenty wykazujące, iż nawet bez znajomości tego efektu współczesna nauka dopuszcza możliwość, aby w szczególnych warunkach całość ciepła mogła zostać zamieniona bezpośrednio na energię elektryczną. Tym zaś osobom które ciągle będą argumentować, iż teoretyczne dopuszczenie możliwości takiej zamiany wcale nie oznacza że będziemy w stanie zrealizować ją technicznie, pragnąłbym rekomendować treść rozdziału S jaki wykazuje, iż komora oscylacyjna już została zrealizowana technicznie. Zgodnie zaś z informacjami świadków obserwujących jej działanie, nie wykazuje ona przy tym zauważalnego nagrzewania dowodząc w ten sposób, iż omawiany tu mechanizm odzyskiwania energii cieplnej iskier faktycznie musiał zostać w niej zrealizowany technicznie.

C6.3.1. Czy w komorze całe ciepło iskier będzie odzyskiwalne

Jedną ze stereotypowych opinii panujących wśród naukowców jest, że zamiana energii cieplnej w jakąkolwiek inną formę energii zawsze musi wypełniać zasadę Carnot'a dotyczącą termodynamicznej sprawności. Wyznawcy tego poglądu automatycznie przenoszą go na komorę oscylacyjną bez rozważenia unikalności warunków występujących w jej wnętrzu. Takie zaś mechaniczne zastosowanie praw termodynamicznych do tej komory jest ogromnym uproszczeniem, przeaczającym następujące niezwykle istotne czynniki:

1. Tzw. "prawa" termodynamiki nie są wcale prawami, a statystycznymi prognozami wynikowego efektu niezliczonych **chaotycznych** wydarzeń.

2. Zachowanie się cząsteczek gazu w obecności silnego pola magnetycznego wykazuje **porządek**, a nie chaos. Stąd też przebieg przemian energetycznych w obrębie komory oscylacyjnej **nie może** być opisany prawami termodynamiki.

3. Nawet bez rozważenia przyszłego sposobu bezpośredniej zamiany ciepła w elektryczność, opartego na wykorzystaniu efektu telekinetycznego, a ściślej jego działania opisanego w podrozdziałach H6.1.3 i K2.4, na naszym obecnym poziomie techniki idealnie sprawne metody konwersji energii cieplnej są też już znane. Dla przykładu, sama zasada przemiany magneto-hydro-dynamicznej, zapewnia idealną sprawność w odzyskiwaniu energii z ciepła (sprawność tą psuje jednak dzisiejsza realizacja techniczna tej zasady). Stąd też, jeśli zamiana energii pozbawiona zostanie chaotycznego charakteru termodynamicznego, jak to stanie się w przypadku komory oscylacyjnej, wtedy owa idealna sprawność odzysku energii może być osiągnięta.

Ponieważ działanie powyższych trzech czynników jest istotne dla komory oscylacyjnej, zaś niektórzy z czytelników mogą nie być dobrze z nimi obznajomieni, poniżej wyjaśniono ich znaczenie w bardziej szczegółowy sposób.

Ad 1. Statystyczny charakter praw termodynamicznych jest doskonale znany od długiego już czasu. James Clerk Maxwell (1831-1879), autor słynnych równań na fale elektromagnetyczne, przedstawił kiedyś dowód bazujący na działaniu tzw. "Démona Maxwell'a". Dowód ten wykazuje, iż w określonych sytuacjach wyjątkowych, ważność praw termodynamicznych przestaje obowiązywać. Poniżej zacytowane zostało co B.M. Stableford napisał o Drugim Prawie Termodynamiki w swojej książce **[1C6.3.1]** "The Mysteries of Modern Science" (London 1977, ISBN 0-7100-8697-0, strona 18):

"Zostało wykazane iż owo prawo termodynamiki jest wynikiem statystycznego zgrupowania dużej liczby wydarzeń raczej niż nienaruszalna zasada jaka rządzi światem żelazną ręką. ... możemy więc dostrzec iż prawa termodynamiki, aczkolwiek dotychczas zawsze działały one w praktyce, w rzeczywistości mogą kiedyś zostać wyeliminowane poprzez rzadką kombinację warunków - nie jest to więc tak bardzo prawo jak statystyczna prognoza."

(W oryginale angielskojęzycznym: "The law of thermodynamics was shown to be a result of the statistical aggregation of a large number of events rather than an inviolable principle ruling the world with an iron hand. ... we can begin to see that although the law of thermodynamics always works out in practice, it could, in fact, be subverted by an extremely unlikely combination of chance happenings - it is not a law so much as a statistical prediction.")

Ad. 2. Jest już doskonale znanym zjawiskiem, iż silne pole magnetyczne zatrzymuje chaotyczne zachowanie się molekuł i organizuje je w uporządkowany wzór. Jednym z zastosowań tego zjawiska jest wytwarzanie pojemnych pamięci komputerowych. Znajduje ono też bezpośrednie wykorzystanie w celach eliminowania energii cieplnej podczas tzw. "chłodzenia magnetycznego" ("magnetic cooling") - patrz książka **[2C6.3.1]** pióra J.L. Threlkeld, "Thermal Environmental Engineering" (Prentice-Hall, Inc., N.J. 1962, strona 152). Stąd też pole magnetyczne jako takie niesie w sobie potencjał spełniania funkcji "Démona Maxwell'a" zdolnego do obalenia ważności praw termodynamicznych. Należy więc się spodziewać, że w obecności potężnych pól magnetycznych, takich jak pola panujące we wnętrzu komory oscylacyjnej, zamiana energii nie będzie podlegała zasadzie Carnot'a.

Ad. 3. Zasada magneto-hydro-dynamicznej zamiany energii nosi właśnie w sobie możliwość pełnego przetransformowania energii cieplnej w elektryczność. Możliwość ta jest doskonale wyrażona poniższym cytowaniem zaczerpniętym z książki **[3C6.3.1]** pióra J.P. Holman "Thermodynamics" (McGraw-Hill, Inc., 1980, ISBN 0-07-029625-1, strona 700) a referującym właśnie do magneto-hydro-dynamicznej zamiany energii:

"Z energetycznego punktu widzenia, ruch siły o określonej wartości przemieszczenia (mechaniczna praca) jest zamieniany na elektryczną pracę (przepływ prądu przeciwko różnicy potencjałów) za pomocą zasady elektromagnetycznej indukcji. Jest to praca-na-pracę zamiana energii i nie jest ona ograniczana przez zasadę Carnot'a." (W oryginale angielskojęzycznym: "From an energy point of view, the movement of force through a displacement (mechanical work) is converted to electrical work (current flow against potential difference) by means of the electromagnetic induction principle. This is a work-work energy conversion and is not limited by the Carnot principle.")

Unikalne warunki panujące w komorze oscylacyjnej eliminują termodynamiczny (chaotyczny) czynnik jaki w normalnych przypadkach zmniejszałby sprawność zachodzących tam procesów, pozwalając na osiągnięcie w niej idealnej efektywności przemian energetycznych.

Rozważania przytoczone w tym podrozdziale wykazują, iż istnieją całkiem realistyczne i dobrze podbudowane przesłanki, jakie sygnalizują możliwość całkowitego odzysku energii traconej we wnętrzu komory oscylacyjnej. Wszystko więc co jest wymagane na obecnym etapie, to abyśmy nie zamykali naszych umysłów przed taką możliwością, ale postarali się zrealizować ją technicznie w tym niezwykłym urządzeniu.

Eliminacja strat energii nie jest jedyną zaletą bezpośredniej zamiany ciepła w elektryczność jaka może zostać zrealizowana w komorze oscylacyjnej. Taka zamiana oddaje nam także do ręki bardzo prostą metodę dostarczania energii do tego urządzenia. Aby zwiększyć jego zasoby energii, wystarczy więc jedynie podgrzewać jego gaz dielektryczny. Takie podgrzanie może zostać uzyskane na drodze cyrkulowania tego gazu przez jakiś wymiennik ciepła, albo też poprzez skierowanie na komorę zwykłego promieniowania słonecznego. Oczywiście istniało też będzie wiele dalszych zastosowań praktycznych tej zamiany. Jednym z ich przykładów zasługujących tutaj na szczególne uwypuklenie jest użycie komór oscylacyjnych magnokraftu do utrzymywania w jego pomieszczeniach stałej i z góry zadanej temperatury. W ten sposób pędniki tego statku wypełniały w nim będą również dodatkową funkcję klimatyzatorów powietrza.

Kombinacja braku strat energetycznych oraz niezależności produkcji pola magnetycznego od ciągłości dostaw energii (patrz podrozdział C6.2) nadaje komorze oscylacyjnej właściwości obecnie charakteryzującej jedynie magnesy stałe. Gdy bowiem urządzenie to raz rozpocznie wytwarzanie swego pola magnetycznego, będzie ono kontynuowało produkcję tego pola przez całe wieki. Jedynym sposobem na zmniejszenie lub wyeliminowanie tego pola będzie wyczerpywanie się energii komory na drodze wykonywania przez nią jakiejś pracy zewnętrznej. Oczywiście wobec braku strat wewnętrznych samo działanie komory nie będzie nigdy w stanie spowodować takiego wyczerpywania się jej zasobów energetycznych.

C6.4. Spożytkowanie niszczycielskiego pola elektrycznego

Wyróżniającą się własnością komory oscylacyjnej jest, iż gromadzi ona na dwóch przeciwległych elektrodach ładunki elektryczne o równej wartości ale za to przeciwnych znakach (tj. taką samą liczbę pozytywów co negatywów). W takich więc warunkach linie sił pola elektrycznego formowanego przez przeciwległe elektrody wiążą się z sobą nawzajem. To z kolei wymusza aby ładunki elektryczne przeskakujące pomiędzy tymi elektrodami wykazywały tendencje do poruszania się po najkrótszych drogach łączących obie elektrody. Stąd w komorze oscylacyjnej tendencja do naturalnego przebiegu ładunków elektrycznych będzie się pokrywała z drogami tych ładunków wymaganymi też i dla prawidłowego działania owego urządzenia. Jako więc wynik końcowy, materiał obudowy komory uwolniony zostanie od niszczycielskiego działania potencjałów elektrycznych, podczas gdy cała siła tych potencjałów będzie skierowana na wytwarzanie pól magnetycznych (zamiast - jak to jest w elektromagnesach - na niszczenie materiałów z jakich to urządzenie zostało zbudowane).

W opisanym powyżej ukierunkowywaniu przepływu energii elektrycznej, komora oscylacyjna całkowicie więc różni się od elektromagnesów. W komorze bowiem owo ukierunkowywanie odbywa się przez wykorzystanie naturalnych mechanizmów przyciągania elektrostatycznego. Natomiast w elektromagnesach wymuszane ono było sztucznie poprzez odpowiednie ukształtowanie warstewek izolacyjnych jakie wymuszały przepływ prądu wzdłuż zwojów przewodnika, podczas gdy istniejące w nim pole elektryczne starało się przepchnąć ten prąd w poprzek zwojów tego samego przewodnika i poprzez warstwę izolacyjną. Stąd należy się spodziewać, iż komora oscylacyjna będzie wykazywała żywotność

nieporównywalnie większą od elektromagnesów, oraz że czas jej używalności nie będzie ograniczany zużyciem elektrycznym materiałów z których została ona wykonana.

Jak niszczące może być takie zużycie elektryczne izolacji elektromagnesu, daje się poznać poprzez analizowanie czasu życia cewek pracujących pod wysokim napięciem. Dobrze znanym ich przykładem jest cewka zapłonowa w samochodach. Izolacja jej zwojów zwykle ulega uszkodzeniu elektrycznemu już po jakichś 7 latach pracy, pomimo iż mechanicznie nie można na niej się dopatrzeć żadnego śladu zużycia. W elektromagnesach niskiego napięcia proces ten jest wolniejszy dlatego czasami może on pozostawać niezauważalnym dla użytkowników. Ale nawet tam sprawa jego wystąpienia jest tylko kwestią czasu.

C6.5. Sterowanie amplifikujące okresu pulsowań pola

Komora oscylacyjna będzie wykazywała bardzo dużą sterowalność. Jak to zostanie szczegółowiej objaśnione w podrozdziale C7.1, kluczem do manipulowania całym jej działaniem będzie okres pulsowań "T" jej pola. Przez zmianę tego okresu przesterowaniu też ulegną wszystkie inne parametry pracy komory. Stąd praktycznie cała działalność sterowania komorą będzie się ograniczała do wpływania na wartość jej okresu pulsowań "T".

Jak łatwo w komorze oscylacyjnej daje się sterować wartością "T" ujawnia równanie (C7) już dyskutowane w podrozdziale C5.6. Na etapie eksploatacji wszystkie czynności sterujące tym urządzeniem można więc ograniczyć jedynie do zmiany wartości jej współczynnika "s". Zmianę tego współczynnika "s" uzyskuje się albo poprzez zmianę ciśnienia gazu wypełniającego komorę albo też poprzez przesterowanie kompozycji tego gazu. Z kolei zmiana "s" wprowadzi zmianę w okresie pulsowań "T" pola komory.

Aby zilustrować istotę tej metody sterowania komorą, warto tu zaznaczyć, że w elektromagnesie jej odpowiednikiem byłaby zmiana parametrów konfiguracyjnych, takich jak oporności obwodów, liczby zwojów oraz geometrycznego wykonania przewodnika. Gdyby te parametry elektromagnesu mogły zostać łatwo zmienione, sterowanie wydatku tego urządzenia posiadałoby przebieg i efekty podobne do tych z komory oscylacyjnej. Tylko więc w takim nierzeczywistym wypadku sterowanie elektromagnesem osiągnięte zostałoby poprzez manipulowanie jego parametrami konfiguracyjnymi, oraz bez konieczności zmiany mocy prądu dostarczanego do jego uzwojeń. Oczywiście w rzeczywistości nie jest możliwym zbudowanie takiego elektromagnesu. To zaś uzmysławia jak nieporównywalnie lepsze jest sterowanie komory w porównaniu ze sterowaniem elektromagnesów.

Efekty takiego sterowania komory są źródłem jej istotnej przewagi nad sposobem sterowania użytym w elektromagnesach. W komorze zmiany stałych gazu dielektrycznego: Ω , ϵ i μ - wywołujące z kolei zmiany w współczynniku "s", nie wymagają manipulowania ilościami energii zawartej w jej polu elektrycznym i polu magnetycznym. Stąd w owym urządzeniu wszystkie czynności sterujące nie wymagają wcale siłowania się z mocą w niej uczajoną. Jako wynik moc urządzeń sterujących nie jest w niej więc zależna od mocy produkowanego pola (tj. słabe urządzenia sterujące są w stanie zmienić efektywnie parametry potężnych pól magnetycznych). Jest to więc wyraźnym przeciwieństwem elektromagnesów w których zmiana pola wymagała zmiany w prądzie elektrycznym tej samej mocy (w ten sposób sterowanie elektromagnesami wymagało zaprzężenia tych samych mocy co wytwarzanie pola).

Oczywiście każda metoda sterowania wprowadza swoje własne ograniczenia i niedogodności. Tak też będzie z systemem opisanym tutaj. Już teraz można przewidzieć, iż będą występowały ograniczenia w zakresie sterowanych wartości, oraz że będzie ono miało wpływ na intensywność nagrzewania powodowaną zmianami w oporności gazu. Jednakże te niedogodności mogą zostać pokonane technicznie, a także są one nieistotne jeśli porównać je z zaletą uczynienia mocy urządzeń sterujących niezależną od mocy sterowanego przez nie pola.

C7. Dodatkowe zalety komory oscylacyjnej ponad elektromagnesami

Wyeliminowanie wszystkich wad wrodzonych elektromagnesów nie jest jedynym osiągnięciem komory oscylacyjnej. Dodatkowo wprowadza ona bowiem kilka zalet operacyjnych jakie nie cechowały jeszcze żadnego innego dotychczas zbudowanego urządzenia. Dokonajmy więc przeglądu najważniejszych z tych jej dodatkowych zalet.

C7.1. Formowanie "kapsuły dwukomorowej" zdolnej do sterowania swym wydatkiem magnetycznym bez zmiany ilości zawartej w niej energii

Dalsze możliwości sterowania wydatkiem komory oscylacyjnej otwierają się kiedy dwa takie sześciennie urządzenia zostają złożone razem w konfigurację zwaną "kapsuła dwukomorowa" - patrz **rysunek C5**. Kapsuła taka składa się z jednej małej komory wewnętrznej "I" (tj. "inner" po angielsku) zawieszona bezdotykowo przy użyciu wyłącznie sił magnetycznych we wnętrzu większej komory zewnętrznej "O" (tj. "outer"). Aby zapewnić bezdotykowe zawiśnięcie komory wewnętrznej bez niebezpieczeństwa iż dotknie ona i uszkodzi komorę większą, długość boku "ao" komory zewnętrznej musi być pierwiastek kwadratowy z "3" razy większa od długości boku "ai" komory wewnętrznej, tj.:

$$a_o = a_i \sqrt{3}$$

(C9)

(tj. bok "ao" równa się bok "ai" pomnożony przez pierwiastek kwadratowy z "3").

Równanie (C9) wyraża wymóg, iż najdłuższa przekątna komory wewnętrznej nie może przekroczyć najmniejszej odległości pomiędzy dwoma równoległymi ściankami komory zewnętrznej.

Obie komory zostają zestawione w ten sposób, że ich osie centralne pokrywają się z osią magnetyczną "m" wynikowej kapsuły. Jednakże ich biegunowość magnetyczna zostaje nawzajem odwrócona, tj. określone bieguny magnetyczne komory wewnętrznej zostają zorientowane dokładnie w odwrotnym kierunku niż te same bieguny komory zewnętrznej. Dla przykładu jeśli komora wewnętrzna skierowywuje swój biegun "S" (south) ku górze - patrz rysunek C5, wtedy komora zewnętrzna skierowywuje ku górze swój biegun "N" (north), i vice versa. Owa przeciwstawna polaryzacja obu komór powoduje, że ich wydatki magnetyczne nawzajem się odejmują (przechwytyują). Efektem tego odejmowania jest, iż wszystkie linie sił pola magnetycznego wytwarzanego przez tą z komór która posiada mniejszy wydatek wcale nie opuszczają kapsuły jako całości, a są przechwytywane przez drugą z komór po czym cyrkulowane przez nią ponownie do komory o mniejszym wydatku. Stąd też pole magnetyczne odprowadzane do otoczenia przez taką kapsułę reprezentuje jedynie algebraiczną różnicę pomiędzy strumieniem magnetycznym produkowanym przez komorę o większym wydatku i strumieniem produkowanym przez komorę o mniejszym wydatku.

W uformowanej w ten sposób kapsule dwukomorowej, odpowiednie sterowanie okresami pulsacji "T" pól magnetycznych wytwarzanych przez składowe komory, umożliwia aby zawartość energetyczna obu komór albo pozostawała na niezmiennym poziomie, albo też przemieszczana była z jednej komory do drugiej. Stąd też obie komory mogą albo wytwarzać taki sam wydatek pola, albo też jedna z nich może produkować wydatek większy od drugiej. Większy z wydatków może być przy tym dostarczany zarówno przez komorę zewnętrzną "O" jak i wewnętrzną "I". Technicznie rzecz biorąc, zrównoważenie lub też przesyłanie energii pomiędzy oboma komorami zależy jedynie od przesunięcia fazowego pomiędzy okresami "To" i "Ti" ich pulsacji. (Z kolei, jak to zostało już wyjaśnione w podrozdziale C6.5, owe okresy pulsacji są sterowane wyłącznie poprzez zmianę współczynników "s" gazów dielektrycznych wypełniających te komory - patrz równanie C7.) Generalnie rzecz biorąc, kiedy obie komory pulsują w zgodzie z sobą, to znaczy kiedy ich przesunięcie fazowe wynosi 0 stopni, 90 stopni, lub wielokrotność 90 stopni, wtedy ich zawartości energetyczne są utrzymywane na tym

samym poziomie bez żadnych zmian. Jednakże kiedy wytworzone zostanie niezerowe przesunięcie fazowe pomiędzy ich pulsacjami, wtedy energia magnetyczna zaczyna przepływać pomiędzy obu komorami. Im owo przesunięcie fazowe bardziej odbiega od 0 stopni lub 90 stopni i stąd bardziej zbliża się do ± 45 stopni, tym więcej energii przepływa z jednej komory do drugiej. Kierunek tego przepływu jest od komory której pulsacje pola uzyskały wyprzedzające przesunięcie fazowe (tj. której okres pulsacji "T" został przyspieszony w stosunku do okresu "T" drugiej z komór) do komory której pulsacje pozostają w tyle.

Aby zilustrować powyższą zasadę przepływu energii pomiędzy obu komorami za pomocą przykładu, wyobraźmy sobie dwoje ludzi na oddzielnych huśtawkach, połączonych z sobą za pośrednictwem gumowego powroza (obie huśtawki reprezentują komory oscylacyjne danej kapsuły, zaś gumowy powróż reprezentuje łączące je pole magnetyczne). Kiedy oboje huśtają się z zerowym przesunięciem fazowym (tj. gdy ich ruchy są identyczne), energia ich oscylacji pozostaje niezmienną. Jednakże kiedy uformują oni przesunięcie fazowe pomiędzy oscylacjami swych huśtawek, wtedy osoba której huśtawka jest w przodzie zacznie pociągać drugą za pośrednictwem gumowego powroza. W ten sposób energia oscylacji będzie przepływać od szybszego huśtacza do wolniejszego.

Kiedy obie komory danej kapsuły dwukomorowej wytwarzają dokładnie takie same wydatki, linie sił pola magnetycznego z komory wewnętrznej "I" formują zamkniętą pętlę z liniami sił pola z komory zewnętrznej "O". Owa pętla z linii sił obu komór jest zamknięta we wnętrzu kapsuły. Stąd też w takim przypadku obie komory mogą produkować pole magnetyczne o niezwykle wysokim wydatku, jednakże pole to w całości "krąży" w obrębie kapsuły i żadna jego część nie wydostaje się na zewnątrz do otoczenia. Pole magnetyczne uwięzione w takiej pętli i hermetycznie zamknięte w obrębie kapsuły jest nazywane "strumieniem krążącym". W ilustracjach z tego rozdziału jest ono oznaczone jako "C" od angielskiego "circulating flux". Strumień krążący wypełnia niezwykle istotną rolę w kapsułach dwukomorowych bowiem wiąże on i zachowuje na potem energię magnetyczną jaka może stanowić źródło energii dla późniejszego jej działania. Stąd strumień krążący w kapsułach dwukomorowych jest odpowiednikiem dla "paliwa" we współczesnych urządzeniach napędowych. Prawdopodobnie w przyszłości budowane więc będą kapsuły których główna i jedyna funkcja polegać będzie właśnie na akumulowaniu energii. Cała energia takich akumulatorów przyszłości zgromadzona będzie w ich strumieniu krążącym, tak że na zewnątrz tych kapsuł nie pojawi się żadne pole magnetyczne.

Jeżeli jednak zawartość energetyczna obu komór kapsuły dwukomorowej jest nierówna (jak to właśnie zilustrowano na rysunku C5), wtedy pole magnetyczne produkowane przez komorę o większym wydatku podzielone zostanie na dwie części: "C" i "R". Część "R", jaką nazywali tu będziemy "strumieniem wynikowym" (po angielsku "resultant flux"), odprowadzana zostaje na zewnątrz kapsuły do otoczenia. Natomiast część "C", nazwana już uprzednio strumieniem krążącym, nadal zamknięta będzie wewnątrz kapsuły. W strumieniu krążącym "C" zawsze więziony będzie cały wydatek komory o mniejszej zawartości energetycznej. Natomiast strumień wynikowy "R" stanowił będzie różnicę algebraiczną wartości wydatków z komory o większej zawartości i komory o mniejszej zawartości energetycznej. Na rysunku C5 większy wydatek jest wytwarzany przez komorę zewnętrzną "O". Stąd to właśnie jej wydatek rozszczepia się na dwa strumienie "C" i "R". Natomiast cały wydatek komory wewnętrznej "I" z tego rysunku jest zaangażowany w strumień krążący "C". Oczywiście w rzeczywistych kapsułach, zależnie od chwilowej potrzeby, możliwe jest takie nasterowanie ich komór, że dowolna z nich może wytwarzać większy wydatek, tj. zarówno zewnętrzna "O" jak i wewnętrzna "I". Stąd również dowolna z tych komór może dostarczać strumienia wynikowego.

Z uwagi na możliwość iż większy z wydatków może być wytwarzany zarówno przez zewnętrzną jak i wewnętrzną komorę, kapsuły dwukomorowe mogą się znajdować w dwóch odmiennych **trybach pracy** jakie tu nazwiemy: (1) z dominacją strumienia WEWNĘTRZNEGO, oraz (2) z dominacją strumienia ZEWNĘTRZNEGO. W trybie dominacji strumienia WEWNĘTRZNEGO, strumień wynikowy "R" wytworzony zostaje właśnie przez komorę wewnętrzną "I", podczas gdy cały wydatek komory zewnętrznej "O" zamknięty zostaje

w strumieniu krążącym "C". Natomiast w trybie dominacji strumienia ZEWNĘTRZNEGO, strumień wynikowy "R" jest produkowany przez komorę zewnętrzną "O", natomiast wydatek komory wewnętrznej "I" jest w całości zamknięty w obrębie strumienia krążącego "C". Wzrokowe różnice w wyglądzie kapsuł pracujących w obu tych trybach zilustrowano w sposób teoretyczny na **rysunku C6**. Różnice te wynikają z faktu, iż pulsujące pole magnetyczne o ogromnych gęstościach jest przeźroczyste jedynie dla obserwatora który patrzy na nie wzdłuż linii sił. Dla obserwatora patrzącego z dowolnego innego kierunku pole takie jest nieprzeźroczyste i wyglądem przypomina czarny dym. Stąd obserwator patrzący na wylot kapsuły dwukomorowej powinien jedynie zobaczyć wnętrze tej komory która produkuje strumień wynikowy rozprzestrzeniający się w kierunku jego oczu. Natomiast zarys drugiej z komór, która produkuje strumień krążący, przyjmował będzie dla niego wygląd optycznej "czarnej dziury" - po więcej szczegółów na temat tego zjawiska patrz podrozdział F10.4.

Teoretyczny wygląd wylotów kapsuł dwukomorowych pierwszej generacji jak je zilustrowano na rysunku C6 wystąpić może jedynie podczas idealnych warunków obserwacyjnych. W rzeczywistości wygląd ten będzie jednak wypaczony najróżniejszymi czynnikami zakłócającymi, najważniejszym z których będzie działanie soczewki magnetycznej objaśnionej na rysunku F32. Stąd jak w rzeczywistości wyglądać będą wyloty tych kapsuł dwukomorowych pokazano na rysunku S5.

Kapsuła dwukomorowa odprowadza do otoczenia jedynie swój strumień wynikowy jaki reprezentuje algebraiczną różnicę z wydatków obu jej komór. Natomiast strumień krążący zawsze pozostaje zamknięty w tej kapsule i nigdy nie osiąga otoczenia. Stąd też owa konfiguracja komór umożliwia niezwykle szybkie i efektywne sterowanie strumieniem wynikowym odprowadzanym do otoczenia. Sterowanie to osiągane zostaje bez żadnej zmiany w całkowitej ilości energii zawartej w kapsule. Energia ta jest jedynie szybko przetrzucana z komory zewnętrznej do wewnętrznej lub vice versa. Powyższe praktycznie oznacza, iż wydatek pola odprowadzanego z kapsuły do otoczenia można łatwo zmienić, podczas gdy ilość energii zawartej w kapsule cały czas pozostaje na tym samym poziomie. Aby zdać sobie sprawę z ogromnych możliwości tego typu sterowania, poniżej opisane zostały najważniejsze stany/atrzybuty pola magnetycznego odprowadzanego do otoczenia przez taką kapsułę.

(1) Całkowite wygaszenie wydatku kapsuły. Jeśli wewnętrzna i zewnętrzna komora zawierają te same ilości energii magnetycznej i stąd produkują takie same strumienie magnetyczne, ich cały wydatek obiega wewnątrz kapsuły dwukomorowej w postaci strumienia krążącego "C" i nic z ich pola nie wydostaje się do otoczenia. Oczywiście w takim przypadku ogromna energia kapsuły nadal pozostaje uwięziona w jej wnętrzu i w każdej chwili może ona zostać przerzucona na jej zewnątrz poprzez proste przesterowanie okresów pulsowań pola "T" w obu komorach składowych.

(2) Płynna lub skokowa zmiana wydatku magnetycznego kapsuły dokonywana w obrębie zakresu od jej minimalnej (tj. zera) do maksymalnej wartości. Taka zmiana w polu wynikowym "R" odprowadzanym z kapsuły do otoczenia wymaga jedynie odpowiedniego przemieszczenia energii magnetycznej z jednej komory do drugiej. Maksymalny wydatek kapsuły uzyskiwany zostaje kiedy jedna z jej komór koncentruje w sobie prawie całą energię, podczas gdy zawartość pozostałej komory jest prawie zerowa.

(3) Wytwarzanie pola magnetycznego jakie zwraca określony biegun magnetyczny ku wybranemu końcu kapsuły. Zależnie która z obu komór (zewnętrzna czy wewnętrzna) osiąga dominujący wydatek, biegunowość strumienia wynikowego "R" będzie odzwierciedlała biegunowość owej dominującej komory.

(4) Niemal natychmiastowe odwrócenie biegunowości wydatku danej kapsuły (tj. zmiana jej bieguna północnego N na południowy S i vice versa). To odwrócenie może zostać zrealizowane za pomocą sterowania kapsułą i poprzez zwykłe przemieszczenie energii magnetycznej pomiędzy komorami (tj. bez potrzeby mechanicznego obrócenia całej kapsuły).

Kolejną zaletą kapsuły dwukomorowej jest jej zdolność do ścisłego sterowania zmian w czasie (tj. krzywej) strumienia wynikowego. Na **rysunku C7** pokazano przykład takiego sterowania, ilustrujący uzyskiwanie strumienia wynikowego jakiego zmiany w czasie posiadają

przebieg tzw. "krzywej dudnienia". Jeśli częstotści pulsowań pola w obu komorach kapsuły różnią się od siebie (np. kiedy komora wewnętrzna wytwarza strumień " F_i " jakiego częstotści pulsowań jest dwukrotnie wyższa od częstotści pulsowań strumienia " F_o " produkowanego przez komorę zewnętrzną), wtedy algebraiczne odejmowanie się obu tych strumieni wytwarza strumień wynikowy " R " (tj. " F_R " na rysunku C7) jakiego zmiany w czasie następują zgodnie z krzywą dudnienia. W ten sposób można uzyskać szeroką gamę zmian strumienia wynikowego " R " poprzez zwykłe przesterowanie częstotści pulsowań pola w komorze zewnętrznej i wewnętrznej (a ściślej ich okresów " T " które powiązane są z częstotściami " f " poprzez równanie (C8): $f=1/T$). Równie prostym staje się więc wyprodukowanie pulsującego strumienia wynikowego przyjmującego kształt dowolnej z krzywych dudnienia, jak strumienia przemienne o dowolnym przebiegu. W każdym z tych przypadków okres pulsowań pola wynikowego może być sterowany z wymaganą dokładnością.

Prawdopodobnie najbardziej jednak istotną z opisywanych tu zalet sterowniczych kapsuły dwukomorowej jest uzyskanie przez nią zdolności zezwalającej na **wytwarzanie stałego pola magnetycznego**. Kiedy bowiem częstotści oscylacji w obu jej komorach są takie same zaś przesunięcie fazowe pomiędzy nimi wynosi zero, wtedy wytwarzane przez nie dwa przeciwstawnie zorientowane strumienie magnetyczne nawzajem eliminują swoje składowe zmienne. Jeśli to zbiega się z identycznością amplitud tych strumieni, wtedy strumień wynikowy " R " staje się nie-oscyłujący (czyli stały w czasie), identyczny w charakterze do pola wytwarzanego przez dzisiejsze magnesy stałe. Zdolność do wytwarzania stałego pola magnetycznego niepomniernie powiększy i tak już ogromny zakres zastosowań tej konfiguracji komór oscylacyjnych.

Z uwagi na bezpośrednią zależność pomiędzy częstotścią " f " i okresem " T " pulsowań pola wyrażoną równaniem (C8), całość opisywanych tutaj czynności sterowania wydatkiem, krzywą i biegunowością strumienia wynikowego uzyskiwana jest poprzez prostą zmianę współczynnika " s " obu komór, jak to już zostało opisane w podrozdziale C6.5.

Powyższe wyjaśnienia ukazują jak łatwe i różnorodne są zdolności sterownicze kapsuły dwukomorowej. Oczywiście, ta łatwość i uniwersalność sterowania będzie posiadała ogromne znaczenie dla przyszłych zastosowań tych zestawów komór. Jest już obecnie możliwe do przewidzenia, że niemal wszystkie systemy napędowe przyszłości będą wykorzystywały kapsuły dwukomorowe zamiast pojedynczych komór. Ze wszystkich urządzeń napędowych opisanych w niniejszej monografii, takie kapsuły będą wykorzystywane w napędzie magnokraftu (patrz opisy z rozdziału F), oraz w magnetycznym napędzie osobistym (patrz opisy z rozdziału E).

C7.1.1. Kapsuły dwukomorowe drugiej i trzeciej generacji

Jak to wyjaśniono w podrozdziale C4.1, oraz podkreślono w podrozdziałach B1 i M6, komory oscylacyjne pierwszej generacji w kształcie kostki sześcienniej budowane będą jedynie w pierwszym okresie rozwoju napędu magnetycznego wehikułów latających. W okresach drugim i trzecim opracowana będzie konstrukcja komór jeszcze bardziej zaawansowanych, które nazywam komorami drugiej i trzeciej generacji. Z tych komór wyższych generacji, m.in. również kapsuły dwukomorowe będą też formowane. Dla zewnętrznego obserwatora kapsuły takie przyjmą inny wygląd od kapsuł pierwszej generacji. Na obecnym etapie naszego rozwoju, my sami nie jesteśmy jeszcze w stanie zbudować żadnej z owych komór oscylacyjnych. Jednak wystawieni jesteśmy na działania cywilizacji szatańskich pasożytów z UFO, które już je zbudowały i użytkują na Ziemi (patrz podrozdział A3 i rozdziały O do W tej monografii). Stąd jest niezwykle istotne abyśmy potrafili odróżnić po ich wyglądzie zewnętrznym te trzy generacje komór oscylacyjnych, oraz odróżnić kapsuły dwukomorowe z nich formowane. Dla umożliwienia takiego odróżnienia, w niniejszym podrozdziale przytoczony zostanie ich pełniejszy opis.

Kapsuły dwukomorowe złożone z komór drugiej i trzeciej generacji pokazane zostały na **rysunku C8**. Ich cechą jest, że podobnie jak kapsuły dwukomorowe pierwszej generacji one również składają się z jednej dużej komory zewnętrznej (O), oraz z jednej małej komory wewnętrznej (I). Owa duża komora zewnętrzna (O) na rysunku C8 zwymiarowana jest średnicą "D" okręgu opisanego na wieloboku ich ścianki czołowej. Natomiast owa mała komora wewnętrzna (I) na rysunku C8 zwymiarowanych średnicą "d" okręgu opisanego na wieloboku ich ścianki czołowej. (Porównaj także rysunki C5 i C8.) W przypadku kapsuł dwukomorowych drugiej generacji zarówno komora zewnętrzna (O), jak i komora wewnętrzna (I) budowane będą w kształcie pręta ośmiobocznego - patrz część (2s) rysunku C8. Stąd jeśli ktoś będzie oglądał je w widoku od czoła (patrz części (2i) i (2o) rysunku C8) wtedy wyraźnie odnotuje że ich ścianki czołowe posiadają kształt ośmioboku równoramienne. Natomiast w przypadku kapsuł dwukomorowych trzeciej generacji zarówno komora zewnętrzna (O), jak i komora wewnętrzna (I) budowane będą w kształcie pręta szesnastobocznego - patrz część (3s) rysunku C8. Przy tak dużej ilości boków, kiedy oglądane będą z odpowiedniego dystansu wtedy obserwatorowi przypomną odcinek niemal okrągłego pręta (patrz części (3s), (3i) i (3o) rysunku C8).

W konstrukcji kapsuł dwukomorowych drugiej i trzeciej generacji spełnionych musi zostać kilka warunków konstrukcyjnych. Warunki te powodują, że wygląd tych kapsuł musi być ściśle określony. Najważniejszym z nich, jakiego znaczenie wyjaśnione jest w podrozdziale C7.1.2, jest tzw. "stopniem upakowania". Stwierdza on, że proporcje wysokości H (h) komór składowych danej kapsuły do średnicy D (d) okręgu opisanego na ich ścianie czołowej muszą być ściśle zdefiniowane. Te proporcje muszą bowiem zapewnić możliwie najwyższą objętość kapsuły przy możliwie najmniejszym zapotrzebowaniu przez nią na cenną przestrzeń statku. Z moich dotychczasowych dociekań wynika, że dla kapsuł drugiej i trzeciej generacji, proporcje te wyniosą $D/H=1$ oraz $d/h=1$. Aby lepiej uświadomić czytelnikowi ich interpretację, proporcje te zilustrowano w części (3s) rysunku C8 poprzez rzadkie zakreskowanie kwadratu o wymiarach D i H oraz podwójnie gęste zakreskowanie kwadratu o wymiarach d i h. Z kolei wymiary komory zewnętrznej (O) i wewnętrznej (I) będą tak dobrane aby wypełnić warunek swobodnego obracania się komory wewnętrznej (I) w komorze zewnętrznej (O) bez powodowania uszkodzenia którejkolwiek z nich. Stąd ich wymiary D i H oraz d i h - patrz części (2o) i (3o) rysunku C8, spełniać również muszą następujące dwa warunki (zauważ że zapożyczony z programowania komputerów symbol "sqrt" oznacza "pierwiastek kwadratowy z argumentu podanego w nawiasie"):

$$A > \sqrt{h^2 + d^2} \text{ oraz } H > \sqrt{h^2 + d^2} \quad (C10)$$

Podobnie jak to ma się z kapsułami dwukomorowymi pierwszej generacji, również kapsuły dwukomorowe drugiej i trzeciej generacji pracować mogą w trybie dominacji strumienia wewnętrznego (patrz części (2i) i (3i) rysunku C8), oraz w trybie dominacji strumienia zewnętrznego (patrz części (2o) i (3o) rysunku C8). Odnotuj, że w magnetycznej konwencji działania, wylot komory której wydatek całkowicie zamykał się będzie w postaci strumienia krążącego przyjmie formę optycznej "czarnej dziury". W zależności więc od tego który z dwóch możliwych trybów dominacji został włączony, wygląd czołowy kapsuły będzie różnił się w sposób charakterystyczny - porównaj części (2i) i (3i) na rysunku C8 z częściami (2o) i (3o) tego samego rysunku.

Oczywiście, rysunek C8 pokazuje teoretyczny wygląd wylotów kapsuł dwukomorowych w magnetycznej konwencji działania. Wygląd ten może jednak wystąpić tylko w niemal idealnych warunkach obserwacyjnych. W rzeczywistości jednak wygląd ten będzie wypaczony faktem wirowania wydatku pędników statku, działaniem soczewki magnetycznej, jonizacją otaczającego statek powietrza, oraz kilkoma innymi czynnikami zniekształcającymi - po więcej szczegółów na temat owych czynników zniekształcających patrz podrozdziały F9 i P2.1.1. Stąd rzeczywisty wygląd tych wylotów będzie nieco odbiegał od teoretycznego. Przykładowo w części D rysunku P19, a także na rysunku P29, pokazany jest rzeczywisty wygląd wylotów

kapsuł dwukomorowych drugiej generacji, kiedy pracują one w konwencji magnetycznej. Z fotografii pokazanych na owych rysunkach wynika, że w rzeczywistych kształtach tych kapsuł wprowadzić z łatwością dopatrzeć się można charakterystycznego dla nich ośmioboku. Jednak poszczególne boki tego ośmioboku są nieco zdeformowane.

Warto też odnotować, że niezależnie od magnetycznej konwencji działania, w której wydatek kapsuł drugiej i trzeciej generacji wytworzy efekty wizualne podobne do formowanych przez wydatek kapsuł pierwszej generacji, kapsuły drugiej generacji pracować też mogą w konwencji telekinetycznej, zaś kapsuły trzeciej generacji w konwencji telekinetycznej i konwencji czasu. Efekty wizualne wzbudzone przez nie w tych odmiennych konwencjach działania będą się różniły od efektów wzbudzanych w konwencji magnetycznej. Przykładowo w konwencji telekinetycznej kapsuły te będą wydzielaly albo białe "światło pochłaniania", albo też zielonkawe "światło wydzielania" - patrz też podrozdziały H6.1 i H6.1.3. Natomiast w konwencji czasu, kapsuła dwukomorowa trzeciej generacji nie tylko wytwarzać będzie w swoim wnętrzu odmienne efekty kolorystyczne, ale także swoiste dla manipulacji czasem efekty ruchowe. Dla przykładu, zewnętrzny obserwator może odebrać jej działanie jakby pokazywane ono było na zwolnionym filmie (patrz zjawisko zawieszanej animacji omówione w podrozdziale M1).

C7.1.2. Stopień wymiarowego upakowania komór oscylacyjnych i jego wpływ na wygląd kapsuł dwukomorowych i konfiguracji krzyżowych

We wszystkich kapsułach dwukomorowych uformowanych z komór oscylacyjnych, parametrem konstrukcyjnym o ogromnej istotności będzie tzw. **stopień wymiarowego upakowania "u"**. Zdefiniowany on może zostać jako: "stopień wymiarowego upakowania (u) kapsuły dwukomorowej, jest to stosunek objętości dwóch proporcjonalnie się kopiujących modeli tej samej komory oscylacyjnej, z których mniejszy o objętości (V_i) całkowicie wpisany zostaje w objętość jakiejś kuli, natomiast większy o objętości (V_o) całkowicie opisany zostaje na powierzchni tej samej kuli", tj.:

$$u = V_i/V_o \quad (C11)$$

W powyższej definicji przez "komory proporcjonalnie się kopiujące" rozumiem dwie komory o identycznych kształtach a jedynie odmiennej skali wymiarowej; czyli takie których wszystkie powierzchnie i krawędzie są do siebie równoległe, zaś wszystkie podobne wymiary liniowe mają się do siebie w dokładnie tej samej proporcji. Warto w tym miejscu podkreślić, że niemal wszystkie komory budowane przez daną cywilizację będą komorami proporcjonalnie się kopiującymi. Powodem tego jest, że aby zbudować komorę której proporcje wymiarowe będą się różniły od komor już budowanych, wszystkie badania rozwojowe i konstrukcyjne musiałyby zostać powtórzone od samego początku. Również ponownego rozpracowania dla niej wymagały też będą metody jej sterowania, urządzenia i czynniki sterujące, a także wszystkie komputery i programy sterujące. To zaś jest przedsięwzięciem bardzo kosztownym i czasochłonnym. Z tego właśnie powodu komora "M" ze standardowej konfiguracji krzyżowej pierwszej generacji pokazanej na rysunku C9 będzie musiała odczekać dosyć długo na swoje zbudowanie. (Jej kształt i wymiary odbiegają wszakże od typowej komory sześciennnej.) Zanim więc ta nieproporcjonalna komora "M" z rysunku C9 zostanie zbudowana, nasza cywilizacja używała będzie konfiguracji krzyżowej pokazanej na rysunku C10, która zawiera wyłącznie komory sześcienne (tj. wyłącznie komory proporcjonalnie się kopiujące).

Geometryczne wymiary danej komory są tym lepiej dobrane, im wartość jej stopnia wymiarowego upakowania jest bliższa $u=1$. Właśnie z uwagi na wartość owego parametru, typowe komory oscylacyjne pierwszej generacji budowane będą w formie sześcianu. (Dla sześcianu dokładna wartość omawianego tutaj stopnia upakowania osiąga $u=0.19245$, albo niemal 20%.) W przypadku komór oscylacyjnych drugiej i trzeciej generacji wartość stopnia

wymiarowego upakowania zależęć będzie od stosunku ich wymiarów "D" do "H" lub "A" do "H". To oznacza, że zależęć on będzie od stosunku wzajemnej odległości "A" dwóch przeciwległych ścianek bocznych komory do jej wysokości "H" (lub od stosunku średnicy "D" okręgu opisanego na czole komory do wysokości "H" tej komory). Stąd dla komór głównych (M) lub zewnętrznych (O) zależęć on będzie od stosunku "D/H" lub stosunku "A/H" ich wymiarów. Natomiast dla komór bocznych (U,V,W,X) lub wewnętrznych (I) od stosunku ich wymiarów "d/h" lub "a/h" - patrz interpretacja owych wymiarów pokazana na rysunku C8.

Szczegółne znaczenie stopnia upakowania "u" dla konstrukcji kapsuł dwukomorowych z komór oscylacyjnych wszystkich generacji wynika ze sposobu na jaki kapsuły te będą użytkowane w magnokraftach. Komory zewnętrzne (O) kapsuł dwukomorowych będą bowiem obracały się bezdotykowo w kulistych obudowach pędników magnokraftu (patrz objaśnienia do rysunku F2). Z kolei we wnętrzu komór zewnętrznych (O) bezdotykowo obracały się będą komory wewnętrzne (I). Stąd obie te komory będą tym przydatniejsze dla magnokraftu, im ich proporcje wymiarowe umożliwią na wymiarowe upakowanie możliwie największej ich objętości przy możliwie najmniejszej średnicy obudowy pędnika do której zostaną one zabudowane. Wszakże gdy upakowanie takie wzrasta, moc pędnika będzie rosła podczas gdy zajmował on będzie coraz mniej miejsca w konstrukcji statku. Stąd więc stopień upakowania "u" komory definiuje sobą wymiarową doskonałość pędnika bazującego na tej komorze.

W przypadku komór oscylacyjnych pierwszej generacji wiadomo że stopień wymiarowego upakowania będzie najwyższy jeśli przyjmą one kształt kostki sześciennej. Stąd typowa komora pierwszej generacji będzie miała kształt sześcianu oboku "A" lub "a", oraz stosunku $A/H=1$ ($a/h=1$). W przypadku komór oscylacyjnych drugiej generacji, analityczne rozwiązanie problemu upakowania nie jest już tak proste i oczywiste. Ja dokonałem więc przybliżonego jego rozwiązania metodami graficznymi. Rozwiązanie to wskazuje, że komora o najwyższej wartości stałej upakowania (u) będzie posiadała stosunek D/H bliski jedności, tj. $D/H=1$. Również w przypadku komór oscylacyjnych trzeciej generacji, rozwiązanie problemu upakowania jakie znalazłem wskazuje, że ich stosunek D/H będzie bliski jedności, tj. $D/H=1$. Oczywiście w tym miejscu zachęcam czytelników aby spróbowali zweryfikować moje ustalenia i rozwiązać w sposób ścisły (analityczny) ten istotny problem konstrukcyjny komór oscylacyjnych drugiej i trzeciej generacji.

Bliski jedności stosunek D/H wymiarów komór oznacza, że figura geometryczna otrzymana po przecięciu typowej komory drugiej i trzeciej generacji płaszczyzną pionową przechodzącą przez jej oś magnetyczną "m" i narożniki dotykające koła o średnicy D, będzie kwadratem o bokach D i H (w kwadracie tym więc $D=H$). Aby wizualnie zilustrować ten kwadrat, zakreśkowałem go w części (3s) rysunku C8 (zakreśkowanie to jest rzadsze w obrębie komory zewnętrznej oraz dwa razy gęstsze w obrębie komory wewnętrznej). Typowe komory oscylacyjne drugiej i trzeciej generacji będą więc posiadały dosyć rzucający się w oczy kształt zewnętrzny, jaki łatwo da się odróżnić od komór sześciennych pierwszej generacji pokazanych na rysunkach C5 i C6. W kształcie tym stosunek szerokości "A" tych komór (lub średnicy "D" okręgu opisanego na ich czole) do ich wysokości "H" będzie przyjmował proporcje takie jak to zilustrowano na rysunku C8.

Podsumowując powyższe rozważania, kapsuły dwukomorowe drugiej i trzeciej generacji, a także wszystkie typowe komory oscylacyjne drugiej i trzeciej generacji, budowane będą w proporcjach wymiarowych jak to zilustrowano na rysunku C8. Z tego względu, po zapoznaniu się z tymi rysunkami ewentualni obserwatorzy tych konfiguracji i komór powinni być w stanie łatwo określić z którą komorą mieli do czynienia. To zaś posiadało będzie swoje następstwo w możliwościach zdefiniowania też i generacji statku kosmicznego w jakim komory te zostały użyte - patrz rozdziały L i M oraz T.

C7.2. Formowanie "konfiguracji krzyżowej"

Kapsuły dwukomorowe nie są jedynymi konfiguracjami w jakie można uformować kilka komór oscylacyjnych w celu zwiększenia sterowalności ich strumienia wynikowego "R". Innym zestawem tych komór, posiadającym nawet jedną więcej możliwość operacyjną niż kapsuły, jest tzw. "konfiguracja krzyżowa". Jej budowa i działanie omówione tutaj zostaną na przykładzie konfiguracji pierwszej generacji pokazanej na **rysunku C9**. (Zauważ, że konfiguracją krzyżową pierwszej generacji nazywana jest konfiguracja uformowana w całości z komór oscylacyjnych pierwszej generacji posiadających kwadratowy przekrój poprzeczny - patrz podrozdział C4.1.) W konfiguracji krzyżowej poszczególne komory zestawione zostały w ten sposób, że jedna z nich, zwana komorą główną (M), otoczona jest przez szereg komór bocznych oznaczonych literami (U), (V), (W) i (X). Komory boczne przylegają do każdej ze ścianek bocznych komory głównej w środku długości tych ścianek. W konfiguracjach krzyżowych pierwszej generacji, których przykład pokazany został na rysunku C9, użyte są cztery komory boczne (U), (V), (W) i (X), ponieważ komora główna posiada tylko cztery boki. Natomiast w konfiguracjach krzyżowych złożonych z ośmiobocznych komór oscylacyjnych drugiej generacji - patrz podrozdział C4.1, komorę główną otaczało będzie aż osiem komór bocznych - tj. jedna komora boczna przylegała będzie do każdego z ośmiu boków komory głównej. Na podobnej zasadzie przy szesnastobocznych komorach trzeciej generacji konfiguracje takie posiadały będą aż szesnaście komór bocznych. Bieguny magnetyczne każdej z komór bocznych zwrócone są w tym samym kierunku, podczas gdy komora główna posiada swoje bieguny magnetyczne ukierunkowane odwrotnie w stosunku do biegunów komór bocznych. Wymiary i objętość poszczególnych komór konfiguracji krzyżowej muszą wypełniać określone warunki konstrukcyjne, których złożona teoria pominięta zostanie w tym miejscu, ale zainteresowani czytelnicy mogą ją znaleźć w podrozdziale F4.1. Mianowicie sumaryczna objętość wszystkich komór bocznych musi być równa objętości komory głównej. Stąd w konfiguracjach krzyżowych pierwszej generacji przekrój poprzeczny każdej z ich pięciu komór, poprowadzony w płaszczyźnie prostopadłej do ich osi magnetycznej, musi być kwadratem o długości boku identycznej jak długość boku pozostałych komór. Objętości oraz wymiary wszystkich komór bocznych (U), (V), (W) i (X) muszą być takie same. Ponieważ objętość komory głównej (M) musi być równa sumie objętości wszystkich czterech komór bocznych, stąd dla konfiguracji krzyżowych pierwszej generacji również i długość komory głównej musi być równa sumie długości komór bocznych (wszakże szerokości komory głównej i komór bocznych są takie same).

Konfiguracja krzyżowa stanowi uproszczony model układu napędowego magnokraftu, którego krótki opis zawarty jest w rozdziale F tej monografii. Również działanie tej konfiguracji imituje działanie napędu magnokraftu. Stąd też reprezentuje ona swoistą miniaturkę magnokraftu. Podobnie jak napęd magnokraftu, jest ona zdolna do wytwarzania nie tylko wszystkich rodzajów pól magnetycznych produkowanych przez kapsuły dwukomorowe, ale także różnych rodzajów pól wirujących których takie pojedyncze kapsuły nie były w stanie wytworzyć. Z tego też względu konfiguracja krzyżowa będzie jedynym zestawem komór możliwym do zastosowania w tzw. magnokrafcie czteropędnikowym (opisanym w rozdziale D) jakiego napęd wymaga właśnie użycia pędników wytwarzających pole wirujące.

W sensie technologii wytwarzania, konfiguracja krzyżowa jest łatwiejsza do zbudowania od kapsuły dwukomorowej. Przyczyną tego faktu jest, iż w kapsule dwukomorowej istnieją techniczne trudności ze sterowaniem komory wewnętrznej, do której wszystkie sygnały sterujące muszą się przedostać poprzez potężne iskry i pole magnetyczne komory zewnętrznej. Trudności te nie występują w konfiguracji krzyżowej w której dostęp z układami sterującymi jest równie łatwy dla każdej z jej komór. Stąd w pierwszym okresie po zbudowaniu komór oscylacyjnych przez spory okres czasu najprawdopodobniej będziemy w stanie formować z nich jedynie konfiguracje krzyżowe. Aczkolwiek więc napęd magnokraftu jest znacznie efektywniejszy jeśli wykorzystuje on kapsuły dwukomorowe w swych pędnikach, owe trudności technologiczne ze sterowaniem takich kapsuł mogą powodować, iż pierwsze dyskoidalne magnokrafty zbudowane na Ziemi będą zawierały właśnie konfiguracje krzyżowe w swoich pędnikach (patrz szczebel rozwojowy 1A w podrozdziale M6). Oczywiście

konfiguracje krzyżowe używane w owych pierwszych dyskoidalnych magnokraftach te nie będą jeszcze konfiguracją z rysunku C9, a specjalnie dla nich opracowaną konfiguracją pokazaną na rysunku C10 i objaśnioną w podrozdziale C7.2.1.

Powyższe również odnosi się do innych cywilizacji już dysponujących magnokraftami. Po tym jaką konfigurację komór wykorzystują one w pędnikach swoich dyskoidalnych wehikułów magnokrafto-podobnych można więc będzie oceniać jak technologicznie zaawansowana jest dana cywilizacja. W pierwszym okresie bowiem po zbudowaniu komór oscylacyjnych najprawdopodobniej każda cywilizacja będzie używała w pędnikach swoich dyskoidalnych magnokraftów nie bardzo tam pasujące konfiguracje krzyżowe, a dopiero potem przeczuci się na trudniejsze technologicznie aczkolwiek bardziej odpowiadające magnokraftom kapsuły dwukomorowe. Podobnie w pierwszym okresie po zbudowaniu komór oscylacyjnych drugiej generacji również najprawdopodobniej cywilizacja ta używała będzie w swoich dyskoidalnych magnokraftach najpierw konfiguracje krzyżowe, a dopiero potem do użycia wprowadzi trudniejsze technologicznie kapsuły dwukomorowe drugiej generacji bazujące na komorach o przekroju ośmiościennym (które zastąpią u niej łatwiejsze do wytwarzania i sterowania komory o przekroju kwadratowym). Dopiero w końcowym okresie swego rozwoju cywilizacja ta przeczuci się na szesnastoboczne komory trzeciej generacji - patrz podrozdział M6.

Zasada sterowania polem wytwarzanym przez konfigurację krzyżową jest prawie identyczna do zasady sterowania tego pola użytej w kapsule dwukomorowej. W podobny więc sposób konfiguracja ta wytwarza dwa strumienie: krążący (C) i wynikowy (R). Tyle tylko iż oba te strumienie cyrkulowane są poprzez otoczenie, zaś jedyna różnica pomiędzy nimi polega na długości drogi jaką ich linie sił zakreślają w swojej cyrkulacji, oraz na liczbie komór przez jakie te linie się domykają (strumień krążący domyka swój obieg przez dwie komory tej samej konfiguracji krzyżowej, natomiast strumień wynikowy tylko przez jedną). Stąd też pole magnetyczne wytwarzane przez konfigurację krzyżową może odznaczać się wszystkimi parametrami jakie opisano już dla kapsuły dwukomorowej. Jedyną dodatkową możliwością konfiguracji krzyżowej nie występującą w kapsule dwukomorowej jest wytwarzanie wirów magnetycznych (tj. pola magnetycznego jakiego linie sił wirują wokół osi magnetycznej "m" tej konfiguracji). Ponieważ wiry takie stanowią niezwykle istotny atrybut napędu magnokraftu i stąd bardziej szczegółowo musiały one być objaśnione w podrozdziale F7, powtórne ich omówienie tutaj zostanie pominięte.

Podobnie jak kapsuła dwukomorowa, również konfiguracja krzyżowa może pracować w dwóch odmiennych trybach pracy jakie nazwiemy "z dominacją strumienia wewnętrznego" (tryb ten pokazany został na rysunku C9 - patrz też rysunek C6 "a") oraz "z dominacją strumienia zewnętrznego" (porównaj części "a" i "b" rysunku C6). W przypadku dominacji strumienia wewnętrznego strumień wynikowy (R) całej konfiguracji produkowany jest przez komorę główną (M). Natomiast w trybie dominacji strumienia zewnętrznego strumień wynikowy (R) całej konfiguracji wytwarzany jest przez komory boczne (U, V, W i X).

Konfiguracja krzyżowa posiada jednak jedną poważną wadę jaka będzie decydowała o jej mniejszym upowszechnieniu od kapsuł dwukomorowych. Wadą tą jest, iż nie daje się w niej całkowicie "wygasić" pola magnetycznego odprowadzanego do otoczenia (chyba że zamiast z komór oscylacyjnych konfiguracja ta zbudowana byłaby z pięciu kapsuł dwukomorowych - co jednak eliminowałoby uzasadnienie dla podejmowania jej budowy ponieważ każda kapsuła dwukomorowa zapewnia niemal te same możliwości sterownicze jak cała konfiguracja krzyżowa). Stąd nawet jeśli cały wydatek konfiguracji krzyżowej cyrkulowany jest w postaci strumienia krążącego "C", ciągle ów strumień krążący wydostaje się na zewnątrz konfiguracji (nie jest więc zamknięty w jej obrębie tak jak to ma miejsce w kapsułach dwukomorowych). Z tego też względu konfiguracje te nie będą się nadawały do wielu zastosowań w których obecność pola magnetycznego jest niewskazana (np. do użytku jako akumulatory energii). Dlatego też, poza krótkim początkowym okresem kiedy to nie potrafili jeszcze będziemy budować kapsuł dwukomorowych, w większości przypadków wykorzystanie konfiguracji krzyżowych ograniczone będzie tylko do urządzeń w których koniecznym jest

wytworzenie wirującego pola magnetycznego (np. do napędu magnokraftu czteropędnikowego opisanego w rozdziale D).

C7.2.1. Prototypowa konfiguracja krzyżowa pierwszej generacji

Konfiguracja krzyżowa posiada tą zaletę nad kapsułą dwukomorową, że jej zbudowanie nie wymaga uprzedniego rozwiązania złożonego problemu sterowania komorą wewnętrzną (I). Wszakże w kapsule dwukomorowej owa wewnętrzna komora (I) pływa swobodnie w środku wnętrza komory zewnętrznej (O). Stąd nie ma do niej bezpośredniego dostępu. Nie daje się też do niej podłączyć żadnych kabli sterujących. Z tych powodów konfiguracje krzyżowe będziemy w stanie budować na długo przedtem zanim opracujemy pierwsze kapsuły dwukomorowe. Wszakże w takich konfiguracjach istnieje dostęp kablowy do każdej z komór. Jednak ze standardową konfiguracją krzyżową pierwszej generacji pokazaną na rysunku C9 i składającą się z czterech sześciennych komór bocznych (U, V, W i X), oraz z jednej wydłużonej komory głównej (M), ciągle związany jest dodatkowy problem konstrukcyjny. Mianowicie komora główna jest w niej cztery razy dłuższa od komór bocznych. Problem ten również jest dosyć trudny do technicznego rozwiązania, bowiem przy stosunku wymiarów odmiennym od typowego sześcianu, od razu zaczynają komplikować się warunki brzegowe panujące na elektrodach komory. Stąd również i przebieg wszelkich zjawisk w jej wnętrzu jest inny. Z kolei te warunki i zjawiska wpływają na układ warunków konstrukcyjnych definiujących stabilną pracę komory, na sposób jej sterowania, na programy sterujące, na pracę komputera sterującego, itd., itp. Praktycznie więc, jeśli niezależnie od komór sześciennych (U, V, W, X) ktoś zechce zbudować również i taką prostopadłościenną komorę główną (M), wtedy niemal cały przebieg badań i rozwoju musi przeprowadzić on od samego początku i to w zakresie znacznie bardziej skomplikowanym. Stąd zapewne zajmie kilka dalszych ładnych lat zanim komora prostopadłościenna (M) będzie gotowa do użytku. Z drugiej zaś strony, niemal natychmiast po opracowaniu pierwszej działającej komory oscylacyjnej, zapewne rząd i społeczeństwo będą naciskały aby badacze ruszyli z rozwojem magnokraftu. W tej sytuacji będą oni zmuszeni do rozpracowania jakiejś sterowalnej konfiguracji komór oscylacyjnych która będzie nadawała się do użycia w pędnikach magnokraftu, a która jednocześnie zawierała będzie jedynie typowe komory w kształcie sześcianów. Konfiguracją tą jest prototypowa konfiguracja krzyżowa pokazana na **rysunku C10**. Używa ona bowiem wyłącznie komór oscylacyjnych w kształcie sześcianu, a więc komór jakie zbudowane będą najwcześniej na naszej planecie.

Rysunek C10 ilustruje ową prototypową konfigurację, pokazując ją z dwóch kierunków. U góry rysunku, tj. w jego części (1s) pokazana jest ona w widoku z boku (side view). Natomiast u dołu rysunku, w jego części (1t) pokazana jest ona w widoku z góry (top view). Dla lepszej informatywności w obu częściach rysunku zaczerniono tworzywo wypełniające które mocuje poszczególne komory i utrzymuje je w wymaganym położeniu i odległości od siebie. Prototypowa konfiguracja krzyżowa składa się z jednej sześciennej komory głównej (M) której bok ma wymiar "A", oraz z ośmiu sześciennych komór bocznych (U, V, W, X) których bok posiada wymiar "a". Przy spełnieniu warunku, wymiarowego że $A=2a$, objętość komory głównej (M) jest równa sumie objętości wszystkich ośmiu komór bocznych (U, V, W, X). Tak dobrana liczba komór prototypowej konfiguracji krzyżowej, ich kształt sześcienny, oraz podane powyżej proporcje wymiarowe powodują, że konfiguracja ta posiada bardzo charakterystyczny wygląd płaskiego dysku o szerokości $G=2A$ i wysokości w swym centrum równej $H=A$ zaś na bokach równej $h=a=(1/2)A$. Stąd bez trudu da się ją wzrokowo odróżnić od standardowej konfiguracji krzyżowej komór oscylacyjnych pokazanej na rysunku C9, jakie zbudowana będzie dopiero w znacznie późniejszym okresie.

Prototypowa konfiguracja krzyżowa z rysunku C10 będzie pierwszą całkowicie sterowalną konfiguracją komór oscylacyjnych jaka początkowo użyta zostanie w wehikułach (typu magnokraft) każdej cywilizacji wkraczającej w swój okres podróży międzygwiazdnych.

Jej główną zaletą operacyjną jest, że bardzo podobne programy sterujące jakie są używane dla sterowania całym magnokraftem typu K3, są też używane do sterowania tej kontroli (wszakże cały magnokraft typu K3 ma taką samą liczbę i takie samo rozłożenie pędników, jak owa kapsuła ma poszczególnych komór). Konfiguracja ta stosowana będzie przez długość okresu przejściowego, tj. począwszy od chwili opracowania przez daną cywilizację jej pierwszego magnokraftu, aż do chwili gdy cywilizacja ta opracuje pierwsze kapsuły dwukomorowe. Odnotowanie jej użycia w pędnikach dyskoidalnego magnokraftu będzie więc zawsze oznaką, że cywilizacja która zbudowała dany statek znajduje się dopiero na samiułku początku swej drogi w kosmos - patrz etap (1A) w klasyfikacji omówionej w podrozdziale M6. Owa prototypowa konfiguracja krzyżowa będzie bowiem wymieniona na przynależną w tym pędniku i bardziej od niej efektywną kapsułę dwukomorową (pokazaną na rysunku C5) natychmiast po tym jak cywilizacja ta zdoła opracować pierwszą niezawodną taką kapsułę.

W Polsce osobą która na własne oczy zaobserwowała taką prototypową konfigurację krzyżową w pędniku głównym dyskoidalnego wehikułu magnokrafto-podobnego (UFO), jest Pan Andrzej Domała - współautor traktatu [3B]. Wobec braku znajomości teorii wyjaśnionych w niniejszej monografii, opisuje on tą konfigurację jako pas złożony z ośmiu kostek sześciennych zawieszony na wirującym słupie stojącym w centrum wehikułu. Ów wirujący słup to po prostu kolumna wirującego pola magnetycznego wytwarzanego przez taką konfigurację i rozprzestrzeniającego się od komory głównej (M) w obu kierunkach. (W momencie oglądania pędnik główny tego statku pracował widać z dominacją strumienia wewnętrznego odprowadzanego z jego komory głównej.) Podczas oglądania z boku taka kolumna wirującego pola widoczna wszakże jest właśnie jako słup czerni - po szczegóły patrz podrozdział F10.4.

C7.2.2. Konfiguracje krzyżowe drugiej generacji

Jak to wyjaśniono w podrozdziale C4.1, oraz podkreślono w podrozdziałach B1, M6 i C7.2.1 komory oscylacyjne pierwszej generacji w kształcie kostki sześcienniej budowane będą jedynie w pierwszym okresie rozwoju napędu magnetycznego wehikułów latających. W okresach drugim i trzecim opracowana będzie konstrukcja komór jeszcze bardziej zaawansowanych, które nazywane tutaj są komorami drugiej i trzeciej generacji. Z komór tych m.in. formowane będą również konfiguracje krzyżowe. Konfiguracje takie dla zewnętrznego obserwatora przyjmą inny wygląd od konfiguracji pierwszej generacji. Na obecnym etapie naszego rozwoju, my sami nie jesteśmy jeszcze w stanie zbudować owych komór oscylacyjnych, jednak wystawieni jesteśmy na działania cywilizacji, które już je zbudowały i użytkują na Ziemi (patrz rozdziały O do W). Stąd jest niezwykle istotne abyśmy potrafili odróżnić te konfiguracje po ich wyglądzie. Dla umożliwienia takiego odróżnienia w niniejszym i następnym podrozdziale przytoczony zostanie ich pełniejszy opis.

Konfiguracje krzyżowe drugiej generacji które złożone są z ośmiobocznych komór drugiej generacji pokazane zostały w częściach (2t) i (2s) **rysunku C11**. Ich cechą jest, że składają się one z jednej ośmiobocznej komory głównej, na rysunku C11 oznaczonej symbolem "M", oraz ośmiu podobnych do niej w kształcie komór bocznych otaczających ją naokoło, na rysunku C11 oznaczonych symbolami "U, V, W i X".

W konstrukcji tych konfiguracji spełnionych musi być kilka warunków konstrukcyjnych i operacyjnych jakie wyjaśnione zostaną w tym podrozdziale. Pierwszym z nich jest wymóg konstrukcyjny i operacyjny aby zestawione ze sobą w tą konfigurację komora główna (M) i komory boczne (U, V, W, X) dotykały się bokami, zaś sąsiadujące ze sobą komory boczne (U, V, W, X) dotykały się narożnikami. Aby go wypełnić wymiary "A" i "D" oraz "a" i "d" poszczególnych komór składowych muszą zostać odpowiednio dobrane. Zgodnie z moimi obliczeniami, wymiary te muszą być tak dobrane, aby $D=1.83d$, oraz $A=1.82a$. To powoduje, że w widoku z góry (top view) ta konfiguracja krzyżowa przyjmie dokładnie wygląd pokazany w części (2t) rysunku C11. Stąd jeśli ktoś będzie oglądał ją w widoku od czoła (patrz część (2t)

rysunku C11) wtedy wyraźnie odnotuje kształt ośmioboku równoramienneho wylotów czołowych ich wszystkich dziewięciu komór składowych. Zauważ, że aby poszczególne komory utrzymywane były w dokładnie przynależnym im wzajemnym położeniu, wolne przestrzenie pomiędzy nimi zalane będą tworzywem wypełniającym, jakie na rysunku C10 i C11 wyróżnione zostało poprzez jego zaczernienie.

Jak to już wyjaśniano w podrozdziale C7.2, wymiary poszczególnych komór oscylacyjnych każdej konfiguracji krzyżowej muszą również tak zostać dobrane, aby wypełnić podstawowy warunek równowagi ich wydatków (patrz podrozdział F4.1). Warunek ten stwierdza, że "objętość przestrzenna V_M " komory głównej (M) musi być równa sumie objętości przestrzennej V_S wszystkich ośmiu komór bocznych (V, V, W, X)". Dla konfiguracji krzyżowej drugiej generacji posiadających osiem jednakowych komór bocznych warunek ten może więc zostać wyrażony matematycznie jako

$$V_M = 8V_S \quad (C12)$$

Teoretycznie rzecz biorąc rozważyć należy dwa sposoby spełnienia warunku (C12), tj. (1) kiedy komora główna przyjmie typowe proporcje wymiarowe (tj. kształt komory proporcjonalnie się kopiującej - patrz podrozdział C7.1.2), natomiast komory boczne otrzymają proporcje wymiarowe wynikające z warunku (C12), lub (2) kiedy komory boczne przyjmą typowe proporcje wymiarowe (tj. kształt komory proporcjonalnie się kopiującej - patrz podrozdział C7.1.2), natomiast komora główna nabierze nietypowych wymiarów wynikających dla niej z konieczności spełnienia warunku (C12). Pierwszy (1) z tych dwóch sposobów może jednak zostać odrzucony z powodów praktycznych, ponieważ dla jego wypełnienia komory boczne musiałyby posiadać wysokość "h" taką że dla nich $h < a$. To zaś z powodu niepożądanych warunków brzegowych pojawiających się wówczas na elektrodach tych komór byłoby wysoce niekorzystne. Stąd ze względów praktycznych konfiguracje krzyżowe drugiej generacji budowane będą na sposób (2), tj. kiedy ich komory boczne (U, V, W, X) przyjmować będą typowe proporcje wymiarowe (tj. ich $h/d=1$) - jak to omówiono w podrozdziale C7.1.2, natomiast komora główna (M) przyjmować będzie proporcje wymiarowe ($H/D=1.28$) wynikające z konieczności spełnienia warunku (C12). Takie skonstruowanie konfiguracji krzyżowych drugiej generacji spowoduje że w widoku z boku (side view) przyjmą one dosyć charakterystyczny wygląd. Wygląd ten pokazano w części (2s) rysunku C11. Jego istotną cechą odróżniającą od np. wyglądu konfiguracji pokazanej na rysunku C10 czy C9 jest, że w ogólnym kształcie konfiguracja drugiej generacji będzie przyjmowała wygląd niemal spłaszczonej kuli o stosunku szerokości G do wysokości H równej $G/H=1.5$. Faktycznie też, jeśli w celu użycia w pędniku wehikułu czteropędnikowego drugiej generacji konfiguracja ta umieszczona będzie w owiewce aerodynamicznej podobnej do owiewek pokazanych na rysunku D1, owiewka ta najprawdopodobniej będzie właśnie w kształcie spłaszczonej kuli albo "dyni", a jedynie w wyjątkowych przypadkach w kształcie całkowicie nieosłoniętej "zębatki" nałożonej na krótkim ośmiobocznym wałku.

Podobnie jak to ma się z konfiguracjami krzyżowymi pierwszej generacji, również konfiguracje drugiej generacji pracować mogą w trybie dominacji strumienia zewnętrznego, oraz w trybie dominacji strumienia wewnętrznego. Ponieważ wyloty komór których wydatek całkowicie zamykał się będzie w postaci strumienia krążącego przyjmą zaczerniony wygląd, w zależności od tego który z owych trybów pracy został włączony wygląd czołowy tej konfiguracji będzie różnił się w sposób charakterystyczny. Warto jednak zauważyć, że podczas trybu pracy z dominacją strumienia wewnętrznego nie wszystkie komory boczne będą równocześnie posiadały całkowicie zaczernione wyloty, a dwie z nich - których wydatek w danym momencie czasu będzie bliski zera - będą wykazywały nieco jaśniejsze wyloty jakie wirowały będą po obwodzie tej konfiguracji.

Warto też odnotować, że niezależnie od magnetycznej konwencji działania, w której wydatek konfiguracji krzyżowych drugiej generacji wytworzy efekty wizualne podobne do formowanych przez wydatek konfiguracji pierwszej generacji, konfiguracje drugiej generacji pracować też mogą w konwencji telekinetycznej. Efekty wizualne wzbudzane przez nie w tej odmiennej konwencji działania będą się różniły od efektów wzbudzanych w konwencji

magnetycznej. Przykładowo w konwencji telekinetycznej konfiguracje te będą wydzielają albo białe światło pochłaniania, albo też zielonkawe światło wydzielania - patrz też podrozdziały H6.1 i H6.1.3.

C7.2.3. Konfiguracje krzyżowe trzeciej generacji

Również i konfiguracje krzyżowe trzeciej generacji przyjmować będą odmienny i bardzo dla nich charakterystyczny wygląd. Wygląd ten pokazano w częściach (3t) i (3s) rysunku C11. Dla umożliwienia czytelnikom jego odróżnienia od wyglądu wszelkich innych konfiguracji komór oscylacyjnych, w niniejszym podrozdziale przytoczony zostanie jego opis i objaśnienie pochodzenia.

Konfiguracje krzyżowe trzeciej generacji złożone są z szesnastobocznych komór trzeciej generacji pokazanych w części (c), rysunku C3, w częściach (3s), (3o) i (3i) rysunku C8 oraz w częściach (3t) i (3s) rysunku C11. Ich cechą jest, że składają się one z jednej szesnastobocznej komory głównej, na rysunku C11 oznaczonej symbolem "M", oraz szesnastu podobnych do niej w przekroju czołowym szesnastobocznych komór bocznych otaczających ją naokoło, na rysunku C11 (3t) oznaczonych symbolami U, V, W i X.

W konstrukcji tych konfiguracji spełnionych musi być kilka warunków konstrukcyjnych i operacyjnych jakie wyjaśnione zostaną w tym podrozdziale. Pierwszym z nich jest wymóg konstrukcyjny i operacyjny aby komora główna (M) i komory boczne (U, V, W, X) dotykały się bokami, zaś sąsiadujące ze sobą komory boczne (U, V, W, X) dotykały się narożnikami. Aby spełnić ten warunek wymiary "D" i "A" oraz "d" i "a" poszczególnych komór składowych muszą więc zostać odpowiednio dobrane. Zgodnie z moimi obliczeniami, wymiary te muszą być tak dobrane, aby $D=4d$ oraz $A=4.143a$. To powoduje, że w widoku z góry (top view) ta konfiguracja krzyżowa przyjmie bardzo charakterystyczny wygląd pokazany w części (3t) rysunku C11. Stąd jeśli ktoś będzie oglądał ją w widoku od czoła (patrz część (3t) rysunku C11) wtedy odnotuje kształt szesnastoboku równoramiennego wylotów czołowych ich wszystkich siedemnastu komór składowych. Kształt ten oglądany w pewnej odległości przy której poszczególne krawędzie boków zaczną się zlewać ze sobą sprawiał będzie wrażenie okrągłego koła o ciągłym obwodzie. (Np. kształt okrągłej wyrzutni pocisków rakietowych.)

Jak to już podkreślano w podrozdziałach C7.2 i C7.2.2, wymiary poszczególnych komór oscylacyjnych każdej konfiguracji krzyżowej muszą tak zostać dobrane, aby wypełnić również podstawowy warunek równowagi ich wydatków (patrz podrozdział F4.1). Przypominając ponownie ten warunek, dla konfiguracji krzyżowych trzeciej generacji będzie on stwierdzał, że "objętość przestrzenna V_M komory głównej (M) musi być równa sumie objętości przestrzennej V_S wszystkich szesnastu komór bocznych (U, V, W, X)". Dla konfiguracji krzyżowej trzeciej generacji po matematycznym wyrażeniu warunek ten przyjmie więc postać, że:

$$V_M = 16V_S \quad (C13)$$

Teoretycznie rzecz biorąc także i dla konfiguracji trzeciej generacji rozważyć również należy dwa sposoby spełnienia warunku (C13), tj. (1) kiedy komora główna przyjmie typowe proporcje wymiarowe (wynikające z warunku maksymalizacji stopnia wymiarowego upakowania - patrz podrozdział C7.1.2), natomiast komory boczne nabiorą proporcje wymiarowe wynikające z warunku (C13), lub (2) kiedy komory boczne przyjmą typowe proporcje wymiarowe, natomiast komora główna nabierze nietypowych wymiarów wynikających ze spełnienia warunku (C13). Drugi (2) z tych sposobów musi jednak zostać odrzucony z powodów praktycznych, ponieważ dla jego wypełnienia komora główna musiałaby posiadać wysokość "H" taką że dla niej $H < A$. To zaś z powodu niepożądanych warunków brzegowych pojawiających się wówczas na elektrodach tej komory byłoby wysoce niekorzystne. Stąd ze względów praktycznych konfiguracje krzyżowe trzeciej generacji budowane będą na sposób (1), tj. kiedy ich komora główna (M) przyjmować będzie typowe proporcje wymiarowe (tj. dla niej $H/D=1$) - jak to omówiono w podrozdziale C7.1.2, natomiast

komory boczne (U, V, W, X) przymować będą proporcje wymiarowe ($h/d=4$) wynikające z konieczności spełnienia warunku (C13). Takie skonструowanie konfiguracji krzyżowych trzeciej generacji powoduje że w widoku z boku (side view) przyjmą one bardzo charakterystyczny i rzucający się w oczy wygląd. Wygląd ten pokazano w części (3s) rysunku C11. (Mi osobiście przypomina on zgrubnie okrągłą wyrzutnię pocisków rakietowych.) Jego istotną cechą odróżniającą od np. wyglądu konfiguracji drugiej generacji pokazanej w części (2s) rysunku C11 jest że w ogólnym kształcie konfiguracja trzeciej generacji będzie przyjmowała wygląd walca w którym wszystkie komory mają tą samą długość. Szerokość G tego walca będzie większa od jego wysokości H, tj. $G=1.42H$, gdzie $H = D$. W walcu tym wyróżnić się da komorę główną (M) posiadającą czoło o dużej średnicy D oraz otaczające ją szesnaście komór bocznych (U, V, W, X) posiadające czoła o czterokrotnie mniejszych od niej średnicach d ($D = 4d$).

Podobnie jak to ma się z konfiguracjami krzyżowymi pierwszej i drugiej generacji, również konfiguracje trzeciej generacji pracować mogą w trybie dominacji strumienia zewnętrznego, oraz w trybie dominacji strumienia wewnętrznego. Ponieważ wyloty komór których wydatek całkowicie zamykał się będzie w postaci strumienia krążącego przyjmą zaczerniony wygląd, zależności od tego który z owych trybów pracy został włączony wygląd czołowy tej konfiguracji będzie różnił się w sposób charakterystyczny. Warto jednak zauważyć, że podczas trybu pracy z dominacją strumienia wewnętrznego nie wszystkie komory boczne będą równocześnie posiadały całkowicie zaczernione wyloty, a cztery z nich - których wydatek w danym momencie czasu będzie bliski zera - będą posiadały jaśniejsze wyloty zaś ich położenia będą wirowały po obwodzie tej konfiguracji.

Warto też odnotować, że niezależnie od magnetycznej konwencji działania, w której wydatek konfiguracji trzeciej generacji wytworzy efekty wizualne podobne do formowanych przez wydatek konfiguracji pierwszej generacji, konfiguracje trzeciej generacji pracować też mogą w konwencji telekinetycznej lub w konwencji czasu. Efekty wizualne wzbudzane przez nie w tych odmiennych konwencjach działania będą się różniły od efektów wzbudzanych w konwencji magnetycznej. Przykładowo w konwencji telekinetycznej konfiguracje te wydzielają będą albo białe światło pochłaniania, albo też zielonkawe światło wydzielania - patrz też podrozdziały H6.1 i H6.1.3. Natomiast w konwencji czasu, konfiguracje krzyżowe trzeciej generacji nie tylko wytwarzać będą w swoim wnętrzu odmienne efekty kolorystyczne, ale także swoiste dla manipulacji czasem efekty ruchowe - np. ich działanie obserwator może odebrać jakby pokazywane ono było na zwolnionym filmie (patrz zjawisko zawieszanej animacji omówione w podrozdziale M1).

C7.3. Nieprzyciąganie przedmiotów ferromagnetycznych

Przywykliśmy już do faktu, iż każde źródło pola magnetycznego przyciąga do siebie różne obiekty ferromagnetyczne. Stąd też jeśli uświadomimy sobie moc pola wytwarzanego przez każdą komorę oscylacyjną natychmiast przychodzi nam do głowy obraz naszych przyszłych noży, widelców i maszynek do golenia ulatujących w powietrzu do sąsiada tylko ponieważ włączył on właśnie zakupioną przez siebie potężną komorę. W tym miejscu nadszedł więc czas na uspokojenie naszych obaw; jednym z bardziej niezwykłych atrybutów kapsuły dwukomorowej i konfiguracji krzyżowej jest, iż wytwarzały one będą pole jakie wcale nie przyciąga przedmiotów ferromagnetycznych. W sensie więc swojego oddziaływania na otoczenie pole to przypominać będzie rodzaj "antygravitacji" opisywanej przez autorów "science fiction", nie zaś zwykłego pola magnetycznego. Niniejszy podrozdział opisuje dlaczego i jak to jest osiąmane.

W obwiedzionej części **rysunku C12** pokazano przybliżony przebieg krzywej pulsowań typowego pola wytwarzanego przez kapsułę dwukomorową. Pole to zwykle przyjmuje postać tzw. "krzywej dudnienia" (po angielsku "beat-type curve") składającej się ze składowej stałej "Fo" oraz składowej zmiennej " ΔF " (porównaj obramowaną część rysunku C12 z rysunkiem

C7). Jest powszechnie wiadomym, że każde źródło stałego pola magnetycznego przyciąga do siebie przedmioty ferromagnetyczne znajdujące się w jego pobliżu. Stąd też jest oczywistym, iż składowa stała " F_0 " pola każdej kapsuły będzie powodowała takie właśnie przyciąganie. Niewiele jednakże osób jest wystarczająco obznajomionych z magnetodynamiką aby także wiedzieć, iż pulsujące pole magnetyczne jakiego przebieg w czasie zmienia się z odpowiednio wysoką częstotliwością " f " indukuje we wszystkich przewodnikach elektryczności tzw. prądy wirowe (eddy currents). Prądy te wytwarzają swoje własne pola magnetyczne, jakie - zgodnie z regułą przekory (kontradykcji) obowiązującą w magnetyźmie - odpychają się od pola które spowodowało ich wytworzenie. W wyniku końcowym, pulsujące pola o odpowiednio wysokich częstotliwościach swych zmian w czasie powodują więc odpychanie przedmiotów ferromagnetycznych. Z tego też powodu, zmienna składowa " ΔF " wydatku pola kapsuły będzie powodować odpychanie takich przedmiotów znajdujących się w jej pobliżu. Siła tego odpychania wzrasta ze wzrostem amplitudy " ΔF " a także i ze wzrostem częstotliwości " f " pulsacji danego pola. Stąd też jeśli tak wysterujemy działanie kapsuły dwukomorowej, że będzie ona zmieniała stosunek " $\Delta F/F_0$ " wytwarzanego przez siebie pola, jednakże w tym samym czasie utrzyma ona jego częstotliwość " f " na niezmiennym poziomie, wtedy mogą wystąpić aż trzy różne rodzaje oddziaływań siłowych pomiędzy tą kapsułą a przedmiotami ferromagnetycznymi z jej otoczenia. Oddziaływania te zilustrowane są na rysunku C12 w formie trzech różnych obszarów wartości przyjmowanych dla danego " f " przez parametry " $\Delta F/F_0$ ", tj.:

(1) Jeśli składowa zmienna " ΔF " pola wytwarzanego przez kapsułę przeważa nad składową stałą " F_0 " tego pola, wtedy sumaryczne oddziaływanie pomiędzy kapsułą i przedmiotami ferromagnetycznymi z otoczenia jest **odpychające**. Na wykresie z rysunku C12 zakres owych oddziaływań odpychających stanowi cały obszar zawarty powyżej "krzywej równowagi".

(2) Jeśli jednak składowa stała " F_0 " dominuje nad składową zmienną " ΔF ", wtedy sumaryczne oddziaływanie pomiędzy daną kapsułą i jej otoczeniem jest **przyciągające**. Na wykresie z rysunku C12 zakres tych oddziaływań przyciągających stanowi całe pole zawarte poniżej "krzywej równowagi".

(3) W końcu jeśli pole wytwarzane przez kapsułę tak wysterować, iż uzyskana jest równowaga pomiędzy składową stałą " F_0 " i składową pulsującą " ΔF ", wtedy przyciąganie całkowicie zneutralizuje odpychanie i vice versa. W takim więc przypadku przedmioty ferromagnetyczne z otoczenia **nie będą przez kapsułę ani przyciągane ani też odpychane**. Na wykresie z rysunku C12 parametry " ΔF ", " F_0 " i " f " pola magnetycznego dla którego nastąpi takie właśnie zneutralizowanie oddziaływań magnetycznych leżą dokładnie na pokazanej tam krzywej. Stąd krzywą tą nazywali będziemy "krzywą równowagi" przyciągających i odpychających oddziaływań magnetycznych.

Krzywa równowagi pomiędzy przyciąganiem i odpychaniem pokazana na rysunku C12 definiuje więc parametry pola magnetycznego jakie w normalnym przypadku wytwarzać będzie każda kapsuła dwukomorowa i konfiguracja krzyżowa. Należy się spodziewać, iż wobec nieszkodliwości tego pola, będzie ono prawie zawsze wytwarzane przez napędy wszystkich wehikułów magnokrafto-podobnych. Pole takie bowiem nie będzie oddziaływać w widoczny sposób na przedmioty ferromagnetyczne zawarte w otoczeniu tych wehikułów, a jednocześnie będzie ono doskonale wypełniało nałożone na nie funkcje napędowe. Z uwagi więc na ową niezwykłą własność tego pola, osoby nieobznajomione z moimi teoriami mogą błędnie posądzać, iż pole to jest innego typu niż magnetyczne, np. że stanowi jakieś nieznane nam jeszcze pole "antygrawitacyjne".

W szczególnych jednakże okolicznościach załoga wehikułów magnokrafto-podobnych będzie mogła przesterować własności wytwarzanego przez siebie pola, włączając wybrany rodzaj oddziaływań na przedmioty z otoczenia. Dla przykładu, jeśli militarnie nastawiony magnokraft będzie ścigał jakiś samolot czy raketę w celu jego przechwycenia, wtedy zmieni on swoje pole z neutralnego na przyciągające. W ten sposób z łatwością będzie on mógł zatrzymać, obezwładnić i uprowadzić ścigany przez siebie obiekt. Podobnie, jeśli taki magnetycznie napędzany wehikuł będzie zamierzał uprowadzić np. samochód wraz z jego

zawartością, wtedy po prostu zawisnie on nad wybranym przez siebie obiektem i zwolna przeniesie go na swój pokład poprzez odpowiednie sterowanie przyciągającymi oddziaływaniami swoich pędników. W takich przypadkach pędniki magnokraftu oprócz swych normalnych funkcji napędowych wypełniały też będą dodatkową funkcję urządzeń zdalnego oddziaływania - patrz etap 1E w podrozdziale M6 (tj. funkcje bardzo podobne jak promień podnoszący opisany w podrozdziale H6.2.1). Oczywiście wystąpią również różne sytuacje kiedy włączenie odpychających oddziaływań okaże się użyteczne. Dla przykładu podczas lotów tych wehikułów w przestrzeni kosmicznej włączane będzie takie odpychanie. W ten sposób wszystkie niebezpieczne obiekty, takie jak meteoryty (w większości przypadków zawierające żelazo), pył kosmiczny, pociski czy satelity, zostaną odepchnięte i odrzucone z drogi owych wehikułów. Także gdy wehikuł taki przelatywał będzie ponad nieznaną czy wrogą sobie planetą, jakiej mieszkańcy będą znani ze strzelania i wysyłania pocisków do wszystkiego czego nie potrafią rozpoznać, wtedy dla własnego bezpieczeństwa załoga takich wehikułów magnetycznych włączy zapewne właśnie owo pole odpychające. Ostrzeżeni nim będą więc mogli śmiać się z pocisków i rakiet lokalnych istot, jakie nie potrafią nawet zbliżyć się do ich technicznie wysoko zaawansowanego wehikułu.

Opisana powyżej możliwość użycia odpowiednio nasterowanych komór oscylacyjnych do formowania "pola podnoszącego" zdolnego do pochwytywania, wynoszenia, oraz manipulowania (np. obracania) wybranych przedmiotów metalowych, z czasem prowadziła będzie do budowania wyspecjalizowanych "urządzeń zdalnego oddziaływania" - patrz etap 1E w podrozdziale M6. Urządzenia te przejmą funkcję współczesnych wind, dźwigów i podajników w przenoszeniu wybranych obiektów z przykładowo Ziemi na pokłady wehikułów typu magnokraft, czy załadowywania części i narzędzi przykładowo ze skrzynek do uchyłków maszyn obróbkowych. W przypadku magnokraftów i komór oscylacyjnych pierwszej generacji zasada działania takich magnetycznych urządzeń zdalnego oddziaływania oparta będzie na opisanej w tym podrozdziale zdolności produkowanego przez nie magnetycznego "pola podnoszącego" do formowania wybranego rodzaju oddziaływania siłowego, zmieniającego się płynnie od odpychania, poprzez działanie neutralne, do przyciągania. Natomiast dla magnokraftów i komór oscylacyjnych drugiej generacji urządzenia te wykorzystywać będą oddziaływania telekinetyczne (po szczegóły patrz opis telekinetycznego "promienia podnoszącego" przytoczony w podrozdziale H6.2.1) zaś w magnokraftach i komorach oscylacyjnych trzeciej generacji oparte one będą na zdalnym manipulowaniu upływem czasu.

W tym miejscu warto odnotować, że magnetyczne urządzenia zdalnego oddziaływania wykorzystujące wyłącznie pole magnetyczne (a nie np. telekinetyczny promień podnoszący) zdolne będą do podnoszenia z ziemi dowolnych obiektów, a nie tylko obiektów ferromagnetycznych. Przykładowo będą one w stanie podjąć z ziemi również osoby i zwierzęta aby wynieść je na pokład magnokraftu. Taka możliwość wynika ze zdolności bardzo silnych pól magnetycznych do wzbudzania magnetyzmu wtórnego w poddanych ich działaniu obiektach nieferromagnetycznych. Jedną z najlepszych jej ilustracji jest krótka notatka pod tytułem "Żaba się uniosła" jaka ukazała się w doniesieniu Gazety Woborczej, wydanie z poniedziałku dnia 14.04.1997 roku i potem powtórzona została w polskojęzycznym kwartalniku UFO, nr 2 (30), wydanie z kwietnia-czerwca 1997 roku, strona 2. Oto treść tej notatki: "Holenderscy i brytyjscy uczeni unieśli żywą żabę w powietrze za pomocą potężnego pola magnetycznego, milion razy silniejszego niż ziemskie pole - doniósł sobotni "New Scientist". Tak silne pole zdeformowało atomy, z których żaba jest złożona, a to sprawiło, że każdy z tych atomów stał się magnesikiem. Żaba została w ten sposób namagnesowana w odwrotnym kierunku, wskutek czego potężne pole przyciągnęło ją i pokonując grawitację uniosło nad ziemię. Uczeni z uniwersytetów w Nottingham (Wielka Brytania) i Nijmegen (Holandia) unosili w ten sposób również rośliny, rybę, konika polnego. - Za pomocą odpowiednio silnego pola można będzie unieść człowieka - twierdzi prof. Peter Main. Pewien uczony amerykański zaproponował nawet, żeby z okazji wejścia w nowe tysiąclecie zademonstrować swobodne unoszenie się człowieka w powietrzu na wysokości 30 metrów. Według uczestników eksperymentu jest to całkiem realne. -Żaba wróciła do swoich

współtowarzyszek, wyglądając na zdrową i szczęśliwą - dodaje prof. Main. PIOC, PAP." Notatka ta bazuje na komunikacie "Frog defies gravity" jaki oryginalnie ukazał się w New Scientist, nr 2077, wydanie z dnia 12 kwietnia 1997 roku, strona 13.

C7.4. Wielowymiarowa transformacja energii

Energia zawarta w komorze oscylacyjnej współistnieje aż w trzech różnych formach, tj. jako: (1) pole elektryczne, (2) pole magnetyczne, oraz (3) ciepło (tj. ciepły gaz dielektryczny wypełniający wnętrze tej komory). Owe trzy formy energii znajdują się w stanie nieustannej transformacji pomiędzy sobą. Ponadto komora jest też w stanie (4) wytwarzać i pochłaniać światło, a także (5) wytwarzać lub konsumować ruch (tj. energię mechaniczną). W końcu komora może też (6) gromadzić i przechowywać ogromne ilości energii przez dowolnie długie okresy czasu (tj. działać jako akumulator energii). Taka sytuacja stwarza unikalną możliwość wykorzystywania komory oscylacyjnej na wiele różnych sposobów (nie zaś tylko jako źródła pola magnetycznego), kiedy to jedna z tych form energii jest do niej dostarczana, zaś inna pozyskiwana, zaś okres czasu upływającego pomiędzy tym dostarczeniem i pobraniem może być dowolnie długi. Następujące formy energii mogą zostać albo dostarczone do, albo też pozyskane z, komory oscylacyjnej: (a) elektryczność przekazywana w formie prądu zmiennego, (b) ciepło zakumulowane w gorącym gazie, (c) energia magnetyczna transformowana za pośrednictwem pulsującego pola magnetycznego, (d) energia mechaniczna przekazywana w formie ruchu komory względem innej komory lub ruchu komory względem pola magnetycznego otoczenia, oraz (e) światło które może zarówno zostać pochłonięte przez strumień krążący komory (patrz opis optycznej "czarnej dziury" z podrozdziału F10.4) lub wytworzone po zamienieniu komory w rodzaj żarówki jarzeniowej (patrz podrozdział F1.3). Zależnie więc od tego która z owych form energii zostanie dostarczona do komory, a która z niej pozyskana, komora oscylacyjna może wypełniać funkcję prawie każdego dotychczas zbudowanego na Ziemi urządzenia do produkowania i/lub transformowania energii. Dla przykładu może ona działać jako: transformator elektryczności, generator elektryczności, silnik elektryczny, silnik spalinowy, ogniwo termiczne, grzejnik, ogniwo fotoelektryczne, reflektor z własną żarówką i baterią wystarczającą na tysiące lat działania, itp. **Tablica C1** zestawia tylko kilka przykładów najszybciej wykorzystywanych zastosowań komory oscylacyjnej, wykorzystujących jej zdolność do wielowymiarowej transformacji energii.

C7.5. Nienawrotne oscylacje - unikalny atrybut komory umożliwiający akumulowanie przez nią nieograniczonych ilości energii

Wróćmy teraz do przykładu huśtawki ilustrującej działanie komory oscylacyjnej. Rozważmy co się z nią stanie w przypadku zwiększania dostarczanej do niej energii kinetycznej. W początkowej fazie, każde dodanie huśtawce energii proporcjonalnie zwiększy amplitudę jej oscylacji. W miarę więc jak nasza dostawa energii się zwiększa, jej ramię będzie wznosiło się coraz to wyżej i wyżej, proporcjonalnie do aktualnie posiadanej przez nią energii. W określonym momencie jednak, bezustanne zwiększanie energii huśtawki spowoduje oparcie się jej ramienia o poziomą belkę do której huśtawka ta została zamocowana, a jaka ogranicza jej wychyły. Dalsze zwiększenie energii spowoduje katastrofę: ramię huśtawki uderzy w ową poziomą belkę i jedno z nich (tj. albo belka albo też ramię) musi zostać zniszczone.

Powyższe ograniczenie konstrukcyjne huśtawki na ilość energii kinetycznej jaką może ona zaabsorbować znalazło już techniczne rozwiązanie. Ktoś bowiem wpadł na pomysł aby zbudować huśtawkę jaka nie posiada górnej poziomej belki. Zamiast tej belki jej ramię zamontowane jest w specjalnej obrotowej osi która umożliwia huśtawce wykonanie pełnych obrotów bez żadnej katastrofy. Jeśli więc zamiast zwykłej, użyjemy huśtawki o takiej

specjalnej konstrukcji, wtedy dalsze dodawanie energii kinetycznej ponad energię jaka poprzednio zniszczyła zwykłą huśtawkę, spowoduje wystąpienie zjawiska które możemy nazwać "nienawrotne oscylacje" (po angielsku "perpetual oscillations"). W huśtawkach o nienawrotnych oscylacjach ich siedzenie zamiast wychylać się do przodu i tyłu, zaczyna zataczać pełne kręgi. Dalsze więc zwiększanie ich energii nie powoduje żadnej katastrofy, a jedynie zwiększa szybkość ich ruchu obiegowego. Oczywiście transformacje energii w takich nienawrotnych oscylacjach ciągle istnieją, jednakże wszystkie występujące w nich zjawiska podlegają już odmiennym prawom niż prawa obowiązujące dla zwyczajnych oscylacji. Najważniejszym atrybutem systemów umożliwiających takie nienawrotne oscylacje jest, iż są one w stanie pochłoniąć więcej energii niż wynosi ich pojemność na energię potencjalną.

Jeśli przeanalizujemy konwencjonalny obwód oscylacyjny z iskrownikiem (Henry'ego), wtedy zauważymy iż jest on podobny do zwyczajnej huśtawki z poziomą belką ograniczającą. Gdy bowiem zaczniemy dodawać do niego energii, wtedy nadejdzie taki moment iż jego kondensator ulegnie przebiciu powodując zniszczenie całego obwodu. Jednakże komora oscylacyjna jest właśnie odpowiednikiem usprawnionej huśtawki - bez owej poprzecznej belki ograniczającej. Umożliwia ona więc uzyskanie nienawrotnych oscylacji.

Mechanizm nienawrotnych oscylacji zachodzących w komorze oscylacyjnej można wyjaśnić jak następuje. Jeśli dodamy w niej dalszej energii do energii już zawartej w jej pęku iskier (przeskakujących powiedzmy z elektrody P_R do P_L - patrz część "c" rysunku C1) wtedy pęk ten nie zaprzestanie przeskoku w chwili gdy przeciwstawna elektroda osiągnie swój potencjał wyładowania "U". Inercja pęku będzie bowiem nadal "pompowała" elektrony z elektrody P_R do P_L , aż cała zawarta w komorze energia przetransformuje się z pola magnetycznego na pole elektryczne. Jednakże w chwili osiągnięcia potencjału "U" przeciwstawna elektroda rozpocznie wyładowanie w odwrotnym kierunku tj. od P_L do P_R , bez oglądania się iż wyładowanie w danym kierunku jeszcze nie zostało zakończone. W ten sposób, jeśli energia komory wejdzie w zakres nienawrotnych oscylacji, w komorze wystąpią przedziały czasu gdy dwa pęki iskier przeskakujące w przeciwstawnych kierunkach zaistnieją równocześnie na tej samej parze elektrod. Pierwszy z tych pęków, nazwijmy go inercyjnym, będzie przeskakiwał z elektrody P_R do P_L , podczas gdy drugi z nich, nazwijmy go aktywnym, będzie przeskakiwał z elektrody P_L do P_R . Taki więc równoczesny przeskok iskier pomiędzy tymi samymi elektrodami w obu kierunkach naraz będzie więc elektromagnetycznym odpowiednikiem dla oscylacji nienawrotnych z omówionych wcześniej huśtawek. Należy tu podkreślić, iż wystąpienie tego unikalnego zjawiska jest tylko możliwe jeśli realizujące go urządzenie potrafi spełnić kilka rygorystycznych warunków konstrukcyjnych, stąd też komora oscylacyjna prawdopodobnie będzie naszym pierwszym i narazie jedynym elektrycznym obwodem drgającym zdolnym do jego wytworzenia.

W tym miejscu możemy sformułować ogólną definicję stwierdzającą, że **"nienawrotne oscylacje mogą być realizowane jedynie w takich systemach oscylujących jakich zdolność do zaabsorbowania energii kinetycznej przekracza ich pojemność na energię potencjalną"**. Taka zdolność jest więc atrybutem czysto konstrukcyjnym. Jest ona uwarunkowana przez określone parametry konstrukcyjne urządzenia oraz przez jego strukturę. W przypadku komory oscylacyjnej będzie ona uwarunkowana liczbą iskier jakie dane urządzenie jest w stanie wytworzyć. Z kolei ta liczba zależy od ilości segmentów (igieł) "p" wydzielonych w każdej elektrodzie. Wyznamy więc teraz minimalną wartość dla "p" wymaganą do zaistnienia w komorze zjawiska nienawrotnych oscylacji.

Jak pamiętamy warunkim tych oscylacji jest, że energia kinetyczna zawarta w polu magnetycznym musi być większa od energii potencjalnej zawartej w polu elektrycznym. Znając równania wyprowadzone dla obwodów oscylacyjnych na ilość ich energii zawartej w obu tych formach, powyższe możemy więc wyrazić w postaci następującej relacji:

$$(\frac{1}{2})L*(U^2/R^2) > (\frac{1}{2})C*U^2$$

Jeśli przekształcimy powyższą relację i zastąpimy otrzymaną w ten sposób kombinację zmiennych przez wartości wyciągnięte z równania (C4), wtedy otrzymamy:

$$p > 2s \quad (C14)$$

Warunek (C14) wyraża liczbę segmentów "p" konieczną do wydzielenia w elektrodach komory oscylacyjnej dla zaistnienia w tym urządzeniu zjawiska nienawrotnych oscylacji.

Jeśli więc potrafimy zbudować i użytkować komorę oscylacyjną w ten sposób, że powyższy warunek zawsze będzie wypełniony, wtedy pojemność tego urządzenia nie będzie wprowadzała żadnego ograniczenia na ilość pochłanianej przez nie energii. Z kolei ta właściwość, w połączeniu z niezależnością komory od ciągłości i efektywności dostawy energii, umożliwi zwiększanie ilości energii zawartej w komorze oscylacyjnej do teoretycznie nieograniczonych wartości.

C7.6. Funkcjonowanie jako pojemny akumulator energii

Zjawisko nienawrotnych oscylacji opisane w poprzednim podrozdziale umożliwia nadanie każdej komorze zdolności do zaabsorbowania teoretycznie nieograniczonych ilości energii. Z kolei ten atrybut, połączony ze zdolnością kapsuły dwukomorowej do całkowitego wygaszenia pola odprowadzanego przez nią do otoczenia (tj. do zamienienia całej swej energii w strumień krążący - patrz podrozdział C7.1), pozwala kapsule dwukomorowej przekształcić się w ogromnie pojemny akumulator. Obliczenia jakie wykonałem dla magnokraftu mogą być przydatne dla zilustrowania poziomu pojemności jaki urządzenie to może zapewnić. I tak kapsuła dwukomorowa o objętości około jednego metra sześciennego, nie będzie wykazywała większych trudności w zakumulowaniu 1.5 TWh (tj. Tera-Watogodziny) energii. Jest to więc odpowiednik dla dwumiesięcznej konsumpcji wszystkich form energii (włączając w to elektryczność, benzynę, gaz ziemny i węgiel) dla całej Nowej Zelandii. Gdyby zaś eksplodować taką jednometrową kapsułę z jej 1.5 TWh zawartością, wtedy zniszczenie wywołane przez jej wybuch byłoby odpowiednikiem eksplozji około jednego miliona ton TNT (tj. 1 megatony TNT).

Pole magnetyczne już obecnie doceniane jest jako doskonały czynnik umożliwiający akumulowanie ogromnych ilości energii. Poprzez użycie przewodników nadprzewodzących, nawet współczesne induktory są w stanie przechowywać ogromne ilości energii przez znaczne okresy czasu. Obecnie istnieje sporo projektów badawczych sprawdzających taką możliwość (np. National University in Canberra, Australia, The University of Texas at Austin, USA). Jednym z bardzo poważnie rozpatrywanych zastosowań komercyjnych było zbudowanie ciężkiego elektromagnesu nadprzewodzącego (cryogenicznego) koło Paryża. Jego zadaniem miało być akumulowanie energii elektrycznej w nocy i późniejsze uwalnianie jej w godzinach szczytu.

Zdolność komory oscylacyjnej do akumulowania ogromnych ilości energii całkowicie rozwiązuje problem jej zaopatrzenia w energię podczas działania. Dla większości bowiem zastosowań wystarczy jeśli zostanie ona w pełni naładowana w chwili wyprodukowania, aby potem służyć bez zasilania aż jej energia jest całkowicie zużyta. Ilości energii jakie daje się zakumulować w tych urządzeniach, zezwalają na ich ciągłe użytkowanie przez setki lat bez żadnej potrzeby dalszego doładowania.

C7.7. Prostota produkcji

Komora oscylacyjna prawdopodobnie będzie stanowiła jedno z najbardziej doskonałych urządzeń jakie ludzka technologia kiedykolwiek wytworzy. Jednakże jego doskonałość wyrażać się będzie głównie w ilości wiedzy wymaganej dla jej prawidłowego

zaprojektowania, a także w ilości badań koniecznych dla właściwego ukształtowania jej działania. Kiedy jednak technologia jej wytwarzania zostanie raz rozpracowana, urządzenie to nie będzie wcale trudne do produkcji seryjnej. Z produkcyjnego punktu widzenia składało się ono będzie bowiem z sześciu prostych ścian, które jedynie będą musiały zostać precyzyjnie zwymiarowane, wykonane i zmontowane. Komora nie posiada przecież części ruchomych, skomplikowanych kształtów, czy złożonych obwodów. Praktycznie więc jeśli wiedza o jej konstrukcji byłaby dostępna, powinniśmy być w stanie wyprodukować ją nie tylko w dzisiejszych czasach, ale nawet tysiące lat temu mając do dyspozycji jedynie ówczesne narzędzia, materiały i technologię naszych przodków (patrz też podrozdział S5).

C8. Wytyczne dla eksperymentów praktycznych nad komorą oscylacyjną

Niniejszym chciałbym zaprosić wszystkich czytelników którzy posiadają ku temu warunki, aby dołożyli własny wkład w rozwój komory oscylacyjnej. Aby ukierunkować ich ewentualne działania twórcze, w tym podrozdziale zestawilem początkowe informacje, jakie powinny pomóc w zastartowaniu ich własnych prac nad tym niezwykłym urządzeniem.

C8.1. Stanowisko badawcze

Zgodnie z doświadczeniami dotychczasowych budowniczych komory oscylacyjnej stanowisko badawcze dla dokonywania eksperymentów nad tym urządzeniem powinno obejmować co najmniej cztery następujące składowe: (1) samą badaną komorę, (2) źródło mocy elektrycznej, (3) magnesy trwałe lub elektromagnes używane do odchyłania drogi iskier w kierunku ścianek komory, oraz (4) urządzenia pomiarowe. Najważniejsze szczegóły każdego z tych elementów podsumowano poniżej.

Ad.(1) **Komora**. Dotychczas przeprowadzone eksperymenty wykazują, iż najoptymalniejszy kształt komory to całkowicie zamknięty sześciian. Dobór wielkości komory jest zadaniem dosyć trudnym i odpowiedzialnym, ponieważ z jednej strony im jest ona większa tym wykonalniejsza technicznie i łatwiej zaobserwować zachodzące w niej procesy, z drugiej jednak strony większe komory wymagają nieproporcjonalnie większych napięć zasilających, więcej elektrod, drogiego materiału, robocizny, itp. Stąd praktycznie jej wielkość nie powinna przekraczać sześciianu o długości boku około 100 mm, zaś prawdopodobnie najbardziej optymalna jest komora o długości boku jedynie około 30 mm. Według dotychczasowego rozeznania w pierwszej fazie eksperymentów najlepszym materiałem na sześć ścianek komory jest pleksiglas (szkło organiczne), ponieważ umożliwia on łatwą obróbkę mechaniczną. W bardziej zaawansowanych modelach konieczne jest jednak użycie wytrzymalszych materiałów, np. szkła kwarcowego. Gaz dielektryczny używany w prototypach budowanych dotychczas stanowi zwykle powietrze pod ciśnieniem otoczenia (rodzaj gazu wypełniającego komorę nabierze istotnego znaczenia na bardziej zaawansowanym etapie badań, tj. podczas dostrajania już działającej komory - patrz etap numer 4 z następnego podrozdziału).

Najbardziej istotnym elementem komory oscylacyjnej są jej **elektrody**. Muszą one być wykonane z jakiegoś materiału neutralnego magnetycznie, sztywnego, wytrzymałego, oraz odpornego na działanie iskier i ozonu. Powinny one być igło-kształtne - jak to już wyjaśniono poprzednio. Im są one cieńsze tym lepiej, jako iż grubsze igły sprzyjają powstawaniu w nich prądów wirowych. Elektrody te powinny być upakowane tak gęsto jak to tylko możliwe bez ich wzajemnego kontaktowania się z sobą. Od gęstości ich rozłożenia zależy wszakże większość cech, parametrów pracy i niepożądanych zjawisk komory, takich przykładowo jak induktancja snopu iskier, pojemność komory, wielkość prądów Halla, oraz wiele innych. Wzajemne rozmieszczenie igieł jest też niezwykle istotne - wszystkie one powinny być w tych samych odległościach od siebie. Dla wypełnienia tego warunku należy je więc ustawiać w układzie

heksagonalnym, tj. takim w którym każda elektroda znajduje się w centrum sześcioboku równoramiennej, jakiego narożniki tworzone są przez elektrody sąsiednie. Najważniejszą częścią elektrod są ich czubki emitujące iskry. Od kształtu i własności tych czubków zależeć będzie powodzenie pierwszych etapów badań eksperymentalnych. Czubki te powinny być zaokrąglone w prawie idealne półkule, jako iż ostre zakończenia powodowałyby prądy ulotowe uniemożliwiające powstanie iskier, zaś płaskie zatępienia wywoływałyby powstawanie niepożądanych zjawisk krawędziowych. Osadzenie igieł w ściankach komory powinno być tak zaprojektowane aby w początkowej fazie eksperymentów umożliwiało łatwą ich wymianę lub regulację wysokości, długości, kształtu, itp.

Ad.(2) Źródło **zasilania** w energię. W pierwszych dwóch etapach budowy komory powinna ona być zasilana prądem zdolnym wytworzyć iskrę co najmniej o długości równej szerokości samej komory. Źródłem takiego prądu może być przykładowo wysokonapięciowy generator impulsów prądu stałego, podobny do tego wykorzystywanego w elektronicznym zapłonie samochodów. Produkowałby on impulsy prądu stałego, jakich zmiana w czasie w przybliżeniu posiadałaby przebieg prostokątny. W dotychczasowych eksperymentach stosowano napięcie impulsów około 300 kV. Należy tu jednak podkreślić, iż po przekroczeniu poza drugi etap budowy komory (jak to zostało opisane w następnym podrozdziale), sposób jej zasilania w energię ulegnie drastycznej zmianie. To z kolei przewartościuje wymagania stawiane zasilaczowi komory. Przykładowo zamiast wysokości wytwarzanego przez niego napięcia oraz kształtu jego impulsów odgrywać w nim rolę zacznie dokładność zsynchronizowania pulsów jego energii z częstotliwością własną komory.

Powinno w tym miejscu zostać dodane, że dobranie lub zbudowanie efektywnego urządzenia do zasilania komory prądem w pierwszych dwóch etapach jej rozwoju może stanowić dosyć trudny i kosztowny problem stanowiska badawczego. Już bowiem wiadomo, że nie każde urządzenie okazuje się przydatne do eksperymentów nad komorą. Dla przykładu iskry wytwarzane przez cewkę Tesli wykazują tendencję do przeskakiwania w niekontrolowanych kierunkach i opierają się próbom wprowadzenia do nich porządku. Z kolei iskry ze wszystkich typów wysokonapięciowych generatorów prądu zmiennego utrzymują swoje kanały jonowe otwarte długo po zaniknięciu iskry, tak że napięcie na elektrodach komory nie jest w stanie już się odbudować.

Uważam wszakże, że przy odpowiednim zaprojektowaniu i właściwym poprowadzeniu programu rozwojowego komory (np. tak jak to zostało uczynione w podrozdziale C8.2) do zasilania komory w pierwszych dwóch etapach jej rozwoju całkowicie powinna wystarczyć zwykła maszyna Wimshurst'a, maszyna Van de Graaff'a, lub nawet połączenie samochodowej cewki zapłonowej (albo induktora wysokonapięciowego) z akumulatorkiem lub baterią. Wszakże gdy w 1845 roku Henry dokonywał eksperymentów nad swoim obwodem oscylacyjnym, jedynym znanym sposobem elektryzowania przedmiotów było ich ręczne pocieranie (maszyna elektrostatyczna Wimshursta została wynaleziona dopiero w 1878 roku) - nie powstrzymało go to jednak przed skompletowaniem rewolucyjnego wynalazku. Oczywiście użycie bardziej złożonych generatorów wysokiego napięcia zwiększy komfort badań. Niemniej nie przybliży ono do osiągnięcia celu, który przecież polega na wypracowaniu efektywnie działających rozwiązań dla samej komory, nie zaś dla jej źródła zasilania.

Ad.(3) **Elektromagnes** (lub układ magnesów stałych) stosowany dla odchyłania iskier. Podczas eksperymentów komora powinna być ustawiana pomiędzy biegunami N i S silnego pola magnetycznego. Pole to przebiegać powinno wzdłuż jej pionowej osi "m", przypierając iskry w kierunku powierzchni ścian bocznych. Owo przypieranie wymusi rotowanie iskier w kierunku zgodnym (lub przeciwnym) do ruchu wskazówek zegara. Bez owego początkowego pola magnetycznego przyłożonego wzdłuż osi "m", iskry nie będą rotowały w uporządkowany sposób naokoło ścianek komory, a raczej przeskakiwały chaotycznie we wszystkich możliwych kierunkach. W chwili gdy efektywność działania komory odpowiednio się zwiększy (patrz koniec etapu 3 w następnym podrozdziale), owa odchyłająca funkcja pola zewnętrznego przejęta zostanie przez pole własne wytwarzane przez daną komorę. Dla wytworzenia wymaganego pola zewnętrznego najkorzystniejszym rozwiązaniem byłby silny

elektromagnes prądu stałego. Prawdopodobnie jednak możliwe też by było użycie w tym celu obwodu magnetycznego uzyskanego przez nałożenie kilku magnesów stałych na odpowiednio zakrzywiony rdzeń ferromagnetyczny którego oba zaostrome końce nacelowane byłyby na oś magnetyczną komory.

Ad.(4) **Urządzenia pomiarowe.** Iskry przeskakujące przez komorę oscylacyjną są bardzo szybkim zjawiskiem jakie jest prawie niemożliwe do dokładnego zaobserwowania gołym okiem i całkowicie opiera się tradycyjnym metodom pomiarowym (począwszy od etapu 3 programu rozwojowego pomiary te nabierają istotnego znaczenia). Z tego powodu stanowisko badawcze powinno obejmować przyrządy pomiarowe dostosowane do obserwacji szybkich przebiegów, dla przykładu oscyloskop, wbudowany aparat lub kamera z wyzwalaniem elektrycznym, megnetometry, itp.

Na zakończenie opisów poszczególnych urządzeń, warto przypomnieć generalną zasadę działalności wynalazczej: "prostota jest kluczem do sukcesu". Odnosi się to nie tylko do urządzeń ale także i do sposobu wprowadzania kolejnych usprawnień do komory jaki powinien podlegać regule "rozkładaj wielkie zadania na szereg małych kroków" (wszakże drogi nawet największych podróżników składały się z wielu pojedynczych kroków). Najoptimalniejsze kompletowanie komory powinno więc nieco przypominać budowę dużego domu z małych cegiełek, którą zawsze zaczyna się od położenia fundamentów, zaś potem prowadzi się systematycznie układając każdą kolejną warstwę znalezisk na warstwie poprzedniej. Patrząc teraz wstecz na dotychczasowy przebieg prac nad rozwojem komory, wszystkich jej początkowych budowniczych wyłożyła właśnie kompleksowość stosowanych przez nich rozwiązań technicznych oraz tendencja do przeskakiwania przez nieistotne ich zdaniem eksperymenty początkowe (np. od razu do etapu 2 "b" lub nawet 3 "b").

C8.2. Etapy, cele i metodyka budowy komory oscylacyjnej

Ponieważ żadne systematyczne badania nad komorą oscylacyjną nie były dotychczas dokonywane, zasadnicza trudność w skompletowaniu tej kryształowej kostki wynika z faktu, iż niemal wszystkie jej szczegóły muszą dopiero zostać odkryte i rozpracowane. Konsekwencją tego jest, że rozwój komory powinien być stopniowy, oraz dokonywany według starannie zaprojektowanego programu (planu). Przy formułowaniu tego planu szczególnie nacisk położono aby realizowany on był według starannie dobranej procedury rozwojowej, jaką nazywam "strategią małych kroków". Podstawowym ogniwem tej procedury jest etap służący rozpracowaniu określonego, ale zawsze tylko jednego, problemu. Z kolei każdy etap daje się rozłożyć na kilka kroków, z których pierwsze zwykle służą modelowemu (tj. dokonywanemu na najprostszym możliwym symulatorze/modelu danego problemu) znalezieniu poszukiwanego rozwiązania, kolejne zaś wypróbowaniu i wdrożeniu tego rozwiązania na rozpracowywanej komorze. W końcu każdy krok, zależnie od użytego sprzętu, napotkanych problemów, oraz uzyskanych wyników, powinien być rozkładany na kilka faz o pojedynczych, jasno zdefiniowanych celach i sposobach ich osiągania.

Po przeanalizowaniu wzajemnych współzależności pomiędzy kolejnymi atrybutami komory oscylacyjnej, jej rozwój daje się rozłożyć na zerowy etap sprawdzający plus osiem prostych etapów rozwojowych. Po owych ośmiu etapach rozwojowych następować może wiele dalszych etapów poszerzających jej funkcjonowanie, dwa z których dla przykładu przytoczono poniżej (patrz etapy 9 i 10). W przypadku takiego rozłożenia, celem każdego kolejnego etapu rozwojowego jest nadanie wynikowej komorze tylko jednej nowej własności użytkowej. Stąd osiągnięcie celu każdego etapu może być dokonane prostymi środkami i z użyciem przejrzystej metodologii badawczej. Owe etapy optymalnego programu budowy komory są jak następuje:

0. **Potwierdzenie poprawności** zasady działania komory. Etap ten nie służy budowie komory, a raczej upewnieniu jej budowniczego oraz osób od których zależy finansowanie lub poparcie jego działań, iż inwestują oni we właściwe urządzenie. Jego celem głównym jest

wykazanie iż generalna zasada działania komory oscylacyjnej nie stoi w sprzeczności z żadnym z praw elektromagnetyzmu i daje się zrealizować na drodze technicznej. Osiągnięcia tego celu można dokonać na wiele sposobów. Przy obecnym stanie rozwoju komory prawdopodobnie najbardziej racjonalny z nich to podjęcie programu realizacyjnego (tj. etapów 1 do 3) oraz następne dodatkowe wykorzystanie dla badań potwierdzających wszelkich zbudowanych w ramach tego programu urządzeń lub modeli komory jakie wytwarzają uporządkowane oscylacje snopów iskier. Cele częściowe w takim wypadku obejmowałyby potwierdzenie że owe snopy iskier: (a) odchylają się zawsze ku tej samej ściance komory (tj. w obecności pola w komorze wykazują one naturalną skłonność do formowania obiegów rotujących wzdłuż jej ścianek), (b) wytwarzają własny strumień magnetyczny podczas takich przeskoków, jaki dodaje się (nie zaś odejmuje!) do strumienia już panującego w komorze, (c) utrzymują się podczas przeskoków jako pęki niezależnych iskier (tj. poszczególne iskry nie łączą się z sobą przed osiągnięciem przeciwstawnych elektrod), (d) wnoszą sobą dodatkową inercję elektryczną (induktancję) do wynikowego obwodu. Należy tu zaznaczyć, że prototypy komory oscylacyjnej dotychczas zbudowane - np. patrz rysunek C13, aczkolwiek pozbawione one były rygorów, systematyki oraz ścisłości ukierunkowania wymaganych dla formalnych badań naukowych, już osiągnęły cel główny tego etapu. Oczywiście wszelkie dalsze ewentualne eksperymenty poszerzające, potwierdzające, weryfikujące, lub formalizujące cel główny lub cele częściowe tego etapu byłyby też jak najmilej widziane (przykładowo szczególnie pożądane byłoby zbudowanie przez kogoś spektakularnego "modelu komory" z rotującymi iskrami, opisanego w kroku "a" etapu 2).

1. Znalezienie podstawowej konfiguracji komory, zdolnej do wytworzenia snopów samo-oscyłujących iskier. Celem głównym tego etapu jest znalezienie takiej konfiguracji elementów komory, jaka wytworzy oscylacyjne wyładowania elektryczne podobne do tych formowanych przez konwencjonalny obwód Henry'ego z iskrownikiem. Aby ułatwić osiągnięcie niniejszego celu głównego, skompletowanie tego etapu należy dokonać na maksymalnie uproszczonym modelu komory jaki składa się tylko z dwóch płytek imitujących ścianki komory, odseparowywanych od siebie płaską przekładką o łatwej do regulowania grubości i utrzymywanych we wzajemnej odległości przez jakieś urządzenie mocujące (np. zwykle imadło ślusarskie). Model ten powinien posiadać tylko jeden obwód oscylacyjny (tj. tylko dwa zestawy elektrod igłowych osadzonych w owych dwóch płytkach ustawionych naprzeciwko siebie). Dla ułatwienia, osiągnięcie celu etapu powinno następować stopniowo, w następujących krokach:

(a) Zbudowanie obwodu inicjującego badania. Celem tego kroku jest praktyczne zainicjowanie badań umożliwiające eksperymentatorowi zapoznanie się z zachowaniem i podstawowymi własnościami obwodów oscylacyjnych z iskrownikiem. Jako pierwszy zbudowany powinien być konwencjonalny obwód drgający z iskrownikiem (tj. obwód Henry'ego z rysunku C1 "a") w którym jednak zamiast pary konwencjonalnych elektrod Henry'ego użyty będzie opisany powyżej uproszczony model komory. W modelu tym wszystkie igły danej ścianki należy zewrzeć z sobą i podłączyć do jednej gałęzi obwodu (np. do induktora i jednej okładziny kondensatora). Stąd w tym kroku zasilanie prądem nastąpi do wszystkich elektrod komory naraz. Po zbudowaniu, tak należy manipulować poszczególnymi parametrami/elementami tego obwodu aby po naładowaniu zmusić go do wytwarzania pęków iskier oscylujących pomiędzy elektrodami komory przez możliwie najdłuższy okres czasu. Im dłuższy czas oscylacji tych iskier, tym łatwiejszy do zaobserwowania będzie potem przebieg eksperymentów z komorą. Należy tu podkreślić że powodzenie tego kroku m.in. zależeć będzie od kształtu i własności czubków elektrod igłowych. Znalezienie więc tutaj najkorzystniejsze ich uformowanie stanowić będzie wkład tego inicjującego eksperymentu przenoszony do następnych etapów badań.

(b) Znalezienie konfiguracji elektrod samo-rozprzestrzeniających iskry. W poprzednim kroku (a) impulsy energii zasilającej dostarczone zostały do wszystkich elektrod komory naraz. Jednakże zastosowane w tym celu rozwiązanie jest nieprzydatne w dalszych badaniach jako że wymagało ono zwarcia z sobą elektrod. Prawidłowo zaprojektowana komora musi więc

przekazywać energię pomiędzy elektrodami na odmienniej zasadzie. Musi ona mianowicie posiadać zdolność do samo-rozprzestrzeniania energii swoich oscylacji. Zdolność ta powodować będzie iż nawet jeśli impuls zasilający dostarczony zostanie jedynie do dwóch jej elektrod igłowych (tj. do jednej elektrody na każdej z obu przeciwstawnych ścianek komory) wskutek wzajemnej indukcji międzyelektrodowej nastąpi rozprzestrzenienie się oscylacji na wszystkie pozostałe elektrody. Niniejszy krok służy nadaniu badanej konfiguracji komory tej właśnie zdolności. Stąd jego celem jest znalezienie takich geometrycznych i konfiguracyjnych parametrów elektrod, jakie spowodują nadanie komorze zdolności do samo-rozprzestrzeniania się iskier. Aby osiągnąć ten cel, w elektrodach modelu komory wypracowanych w efekcie kroku (a) należy teraz dokonywać dalszych modyfikacji geometrycznych i konfiguracyjnych. Kluczem do sukcesu będzie tu czynna długość elektrod (należy pamiętać że długość całkowita elektrod może być zwiększana nie tylko w obrębie komory, ale także poza komorą - na jej zewnątrz). Przykładowo należy więc zwiększać stosunek wysokości lub całkowitej długości tych elektrod do wielkości przerwy międzyelektrodowej, stosunek długości elektrod do ich wzajemnej odległości od siebie, itp. Po osiągnięciu atrybutu samo-rozprzestrzeniania się iskier, cel tego kroku zostanie osiągnięty. Po skompletowaniu tego kroku, wynikowy optymalny obwód oscylacyjny należy zachować, ponieważ będzie on jeszcze przydatny w dalszych etapach badań (patrz etapy 2 "a" i 3 "a").

(c) Zastąpienie induktora roboczego z konwencjonalnego obwodu Henry'ego przez indukcyjność pęku iskier. Celem tego kroku jest znalezienie parametrów konstrukcyjnych i geometrycznych elektrod, koniecznych dla wytworzenia wymaganej indukcyjności obwodu oscylacyjnego wyłącznie przez snopy iskier przeskakujących w komorze. Osiągnięcie tego celu polega na takim manipulowaniu kształtem i właściwościami elektrod w badanym modelu komory (np. poprzez dodanie nieprzewodzących kulek na ich czubkach), ich długością aktywną, średnicą, liczbą, wzajemnymi odstępami, oraz sposobem rozstawienia, aby uzyskane zostało zamierzone zwiększenie indukcyjności snopów iskier. Docelowa indukcyjność wymagana dla skompletowania tego kroku, musi umożliwiać wynikowemu obwodowi na wytwarzanie samo-oscyłujących iskier nawet jeśli induktor zostanie od niego całkowicie odłączony. Po osiągnięciu celu tego kroku, induktor roboczy należy wyeliminować już na stałe z dalszych prototypów komory, zaś używać jedynie właśnie wypracowanej konfiguracji elektrod przy których to iskry, a nie induktor, dostarczały będą obwodowi wymaganej przez niego indukcyjności własnej.

Oczywiście wyeliminowanie induktora roboczego nie oznacza wcale że wynikowy obwód nie powinien zawierać żadnej cewki. Może bowiem się okazać że istnieje konieczność pozostawienia niewielkiej cewki wypełniającej funkcje sterujące (ale nie funkcje robocze - tj. nie funkcje dostarczania obwodowi wymaganej indukcyjności) jaka łączyłaby centralne elektrody/igły obu przeciwległych płyt komory (sukces w zrealizowaniu etapu 2 może nawet zależeć od istnienia takich cewek sterujących). To samo stosuje się też do eliminowania kondensatora roboczego w następnym kroku (d).

(d) Zastąpienie kondensatora roboczego z obwodu Henry'ego przez pojemność własną komory. Celem tego kroku jest wymagane powiększenie pojemności komory poprzez zmianę jej parametrów konfiguracyjnych. W celu skompletowania tego kroku, model komory uzyskany w efekcie kroku (c) należy teraz tak przekształcić poprzez manipulowanie jego parametrami posiadającymi wpływ na pojemność własną, aby po wyeliminowaniu z obwodu także zewnętrznego kondensatora ciągle wytwarzał on samo-oscyłujące pęki iskier. Przykładowe wielkości jakie należy zmieniać aby osiągnąć ten cel to: stosunek przerwy międzyelektrodowej (tj. odległości pomiędzy czubkami elektrod obu przeciwległych ścianek) do odległości poszczególnych elektrod od siebie, stosunek wysokości elektrod do ich wzajemnej odległości, stosunek odsłoniętych do zaizolowanych części elektrod, całkowita liczba elektrod, kształt elektrod, itp. Po dobraniu parametrów jakie umożliwią wytworzenie samo-oscyłujących snopów iskier już po odłączeniu zewnętrznego kondensatora, podstawowa konfiguracja komory oscylacyjnej będzie znaleziona. Konfiguracja ta po naładowaniu elektrycznym dwóch jej elektrod będzie wytwarzała snopy oscylujących iskier (tj. dostarczała "oscylacyjnej

odpowiedzi") nie zawierając przy tym ani zewnętrznego induktora roboczego ani też zewnętrznego kondensatora roboczego.

2. **Samo-regulacja przesunięcia fazowego** pomiędzy dwoma pękami iskier. Kolejnym etapem budowy komory powinno być złożenie razem dwóch końcowych obwodów komory uzyskanych w efekcie zrealizowania etapu 1. Niestety, po złożeniu obwody te - zamiast przeskoków uporządkowanych z wymaganym przesunięciem fazowym wynoszącym 90 stopni, będą raczej wykazywały tendencję do nieskoordynowanych przeskoków iskier. Stąd też celem tego etapu jest wypracowanie takiej konfiguracji (kształtu) komory oraz jej elektrycznych sprzężeń wewnętrznych, aby samo-regulowała i samo-utrzymywała ona 90-stopniowe przesunięcie fazowe pomiędzy oscylacjami zachodzącymi w jej obu obwodach składowych. Droga do osiągnięcia tego celu wiedzie przez wprowadzenie do konstrukcji komory różnorodnych dodatkowych elementów lub zmian, takich jak przykładowo: zaizolowane płyty dołączone do każdej elektrody jakie bezdotykowo zachodzą na elektrody następnych ścianek formując w ten sposób pomiędzy nimi dodatkową pojemność wymuszającą wymagane przesunięcia fazowe (patrz rysunek S7); wybrania w elektrodach podobne do tych formujących stacjonarne fale w kuchenkach mikrofalowych; cewki podobne do cewek rozruchowych stosowanych w jednofazowych silnikach elektrycznych; itp. Dla ułatwienia, podobnie jak w etapach poprzednich, osiągnięcie celu tego etapu powinno nastąpić w dwóch krokach:

(a) Modelowe wypracowanie efektywnego systemu samo-regulacji 90-stopniowego przesunięcia fazowego w dwóch niezależnych obwodach oscylujących. Celem tego kroku byłoby znalezienie, przy wykorzystaniu prostych w budowie i działaniu obwodów Henry'ego, owego efektywnego systemu samo-regulującego.

Dla zrealizowania tego celu, dwa konwencjonalne obwody Henry'ego, poprzednio opisane w kroku 1 (b) lub nawet 1 (a) powinny zostać złożone razem w celu uformowania "modelu komory". W modelu tym dwa układy elektrod zamontowanych na bocznych ściankach komory sześcienniej użyte byłyby jako przerwy iskrowe obu konwencjonalnych obwodów Henry'ego. Obwody te oscylowałyby ze wzajemnym przesunięciem fazowym wynoszącym 90 stopni. Stąd przy obecności zewnętrznego pola magnetycznego formowałyby one w modelu komory snopy iskier rotujących po obwodzie kwadratu. Prosty system jaki byłby w stanie efektywnie utrzymywać wymagane 90-stopniowe przesunięcie fazowe w oscylacjach obu tych obwodów, najprawdopodobniej dostarczyłby zasady dla właśnie poszukiwanego systemu samo-regulującego, nadającego się do adaptacji w wynikowej konfiguracji komory. Warto tu dodać, iż powyższy "model komory" powinien już wytwarzać niewielkie pole magnetyczne, stąd sam w sobie byłby on sporym osiągnięciem naukowym i wynalazczym, nadającym się do opublikowania i popularyzacji technicznej.

(b) Praktyczne wdrożenie właśnie wyracowanego systemu samo-regulacji wymaganego przesunięcia fazowego. Celem tego kroku byłoby takie adaptowanie systemu wyracowanego w kroku (a) aby działał on równie efektywnie w aktualnej konfiguracji komory. Właściwie adaptowany system powinien dawać pęki iskier jakie przeskakiwałyby z wymaganym przesunięciem fazowym wynoszącym 90 stopni pomiędzy dwoma parami przeciwległych ścian komory, jeśli zasilanie w energię nastąpiłoby tylko do jednego z jej obwodów oscylujących (tj. drugi z obwodów powinien samoczynnie zaabsorbować wymaganą przez siebie energię z tego pierwszego obwodu).

3. Zmuszenie komory do zaabsorbowania ilości energii jaka wystarczy na **wytworzenie użytecznego pola magnetycznego**. Celem tego etapu jest znalezienie i zrealizowanie sposobu (techniki) dowolnego zwiększania poziomu energii zawartej w komorze na drodze zasilania jej impulsami magnetycznymi (nie zaś impulsami elektrycznymi jak w etapach poprzednich). Z kolei zwiększenie poziomu tej energii: (a) wprowadzi możliwość nieograniczonego wydłużania czasu trwania wyładowań oscylacyjnych komory, (b) pozwoli na wyeliminowanie zewnętrznego źródła pola magnetycznego jakie przy krótkich impulsach działania komory niezbędne było dla wymuszenia uporządkowanego obiegu iskier, oraz (c) umożliwi komorze wytworzenie własnego pola magnetycznego odprowadzanego z niej do

otoczenia. Główna zasada na jakiej opiera się osiągnięcie celu tego etapu polega na odwróceniu kierunku transformacji energii w komorze (tj. zamiast jak poprzednio tylko transformować prąd zasilania na własne pole magnetyczne, teraz będzie ona transformować pole zasilania na prąd własny, potem zaś prąd własny na własne pole magnetyczne). Droga do zrealizowania celu tego etapu wiedzie przez: (1) znalezienie warunków najefektywniejszego przepływu energii do komory (np. znalezienie punktu w cyklu oscylacji własnych komory jaki najoptymalniej nadaje się dla dostawy impulsu zasilającego, wzajemnego przesunięcia fazowego pomiędzy ciągiem impulsów zasilających i drganiami własnymi komory, najefektywniejszej różnicy amplitud, itp. - patrz podrozdział C7.1), (2) znalezienie sposobu na automatyczne (elektroniczne) wykrywanie wybranego przez nas punktu w cyklu oscylacji własnych komory (tj. punktu który najoptymalniej nadaje się na dostarczenie komorze impulsu zasilającego), (3) znalezienie techniki zsynchronizowanego wyzwalania dostawy impulsów energii z zewnętrznego źródła, następującego dokładnie w wybranym przez nas punkcie cyklu oscylacji własnych komory, (4) zbudowanie urządzenia sterującego jakie efektywnie zrealizuje tą technikę u używanego przez nas zestawu komory i jej źródła zasilania. Jeśli cel tego etapu zostanie osiągnięty, komora będzie w stanie zaabsorbować i przetransformować na pole magnetyczne każdą wymaganą ilość energii. Z kolei energia ta zezwoli komorze na wytwarzanie pola magnetycznego o wymaganym natężeniu oraz na jej nieprzerwaną pracę przez dowolnie długi okres czasu. Co za tym idzie umożliwi ona praktyczne użytkowanie wytwarzanego przez tą komorę pola magnetycznego. Po zrealizowaniu więc tego etapu, prototyp komory zacznie nadawać się do pierwszych zastosowań praktycznych. Najważniejsze kroki realizacyjne są tu jak następuje:

(a) Modelowe wyznaczenie warunków najefektywniejszego przepływu energii do komory. Celem tego kroku byłoby wyznaczenie: (1) wartości różnicy pomiędzy częstotliwością zewnętrznego źródła zasilającego, a częstotliwością własną/rezonansową komory, jaka spowoduje iż komora zaabsorbuje z zewnętrznego źródła i przechowa wymaganą ilość energii; (2) optymalnego przesunięcia fazowego pomiędzy pulsowaniami obu tych elementów; (3) technicznego sposobu "dostrajania" się jednego z elementów (tj. komory albo zewnętrznego źródła energii) do wymaganej częstotliwości i przesunięcia fazowego.

Dla ułatwienia, realizację tego kroku należy dokonać na uproszczonym "modelu kapsuły dwukomorowej" lub "modelu transformatora". Model ten uzyskany byłby poprzez magnetyczne sprzęgnięcie z sobą dwóch konwencjonalnych obwodów drgających. Sprzęgnięcie to nastąpiłoby na drodze magnetycznej za pośrednictwem ich induktorów. Możliwe są przy tym dwa rozwiązania, jakie z uwagi na charakter przyszłego zastosowania oba muszą bazować na induktorach o rdzeniu powietrznym (tj. cewkach posiadających prześwit przez swoje centrum). W pierwszym rozwiązaniu użyty byłby "model kapsuły dwukomorowej" uzyskiwany poprzez wstawienie mniejszego aktywnego induktora powietrznego do wnętrza drugiego pasywnego (podczas praktycznego wdrażania tego modelu cewka zasilająca komorę w energię wstawiana byłaby do wnętrza tej komory). W drugim rozwiązaniu oba induktory w przybliżeniu tej samej wielkości ustawiane byłyby obok siebie jak uzwojenia pierwotne i wtórne zwykłego transformatora (podczas wdrażania tego modelu cewka zasilająca ustawiana byłaby na przedłużeniu osi magnetycznej komory). Po takim magnetycznym sprzęgnięciu, jeden z tych obwodów (aktywny) zasilany byłby w energię drugi z obwodów (pasywny) jakim mógłby być konwencjonalny obwód Henry'ego opracowany w efekcie kroku 1 (b) lub nawet kroku 1 (a). W ten sposób zdefiniowane mogłyby zostać warunki (przesunięcie fazowe lub różnica częstotliwości pulsowań) przy jakich przepływ energii od obwodu aktywnego do pasywnego jest najefektywniejszy. Zaletą użycia takiego uproszczonego modelu jest że obwodem aktywnym może wtedy zostać praktycznie dowolny obwód umożliwiający regulowanie częstotliwości swych drgań w zakresie obejmującym częstotliwość własną obwodu pasywnego (nie byłoby więc konieczne budowanie obwodu wysokonapięciowego). Można by więc w tym celu wykorzystać gotowe obwody oscylacyjne, np. obwody dostrajające ze starych radioodbiorników. Ponadto po zakończeniu badań obwodów

badawczy i urządzenie aktywne mogą zostać adaptowane niemalże bez zmian technicznych do zasilania w energię opracowywanej właśnie komory oscylacyjnej.

(b) Modelowe wypróbowanie znalezionej komory. Celem tego etapu byłoby sprawdzenie w działaniu najprostszego urządzenia jakie zrealizowałyby wyznaczone poprzednio warunki na optymalniejszego przekazywania energii do komory. Dla jego osiągnięcia, zbudować należy prototypowy system automatycznie przekazujący energię do obwodu pasywnego. Użyta metodyka realizacji byłaby podobna jak w kroku 3 "a", tyle tylko iż zamiast służyć znalezieniu najoptymalniejszych warunków i sposobów dostawy energii, urządzenie to starałoby się uczynić z nich możliwie najlepszy użytek.

(c) Praktyczne wdrożenie na komorze wyznaczonych warunków i urządzenia gwarantujących efektywny przepływ energii od zasilacza magnetycznego do rozpracowywanej komory oscylacyjnej. Aby dokonać takiego wdrożenia aż trzy współpracujące z sobą urządzenia muszą zostać zestawione w jeden efektywnie kooperujący zestaw. Są to: (1) komora której elementy (np. czujniki, cewki) oraz żywotność umożliwiają uzupełnianie jej zasobów energii na wypracowanej przez nas drodze, (2) zewnętrzne źródło pulsującej energii magnetycznej (zasilacz), jakie będzie współpracowało z tą komorą w sposób wymagany przez daną technikę, zaopatrując ją efektywnie w energię konieczną do jej ciągłej pracy, oraz (3) urządzenie sterujące jakie będzie koordynowało uzupełnianie zasobów komory przez to zewnętrzne źródło energii, umożliwiając w ten sposób nieprzerwaną pracę całego zestawu.

Należy tu podkreślić że po zakończeniu tego etapu dalsze zasilanie komory w energię odbywać się już będzie za pomocą wypracowanego tutaj systemu generacji impulsów magnetycznych, zaś zasilacz wysokonapięciowy przestanie być potrzebny. W zasilaniu takim komora oscylacyjna będzie teraz pełnić funkcję jakby uzwojenia wtórnego transformatora, którego uzwojeniem pierwotnym jest cewka zasilacza wytwarzająca odpowiednio zsynchronizowane impulsy pola magnetycznego.

4. **Sterowanie okresem pulsowań** komory (z kolei ów sterowany okres pulsowań umożliwia zmianę wszelkich pozostałych parametrów pracy komory - patrz podrozdziały C7.1, C6.5 i C5.6). Celem tego etapu jest poznanie sposobu w jaki można sterować okresem pulsowań (częstotliwością) pola komory poprzez odpowiednie dobieranie i szybką zmianę parametrów zawartego w niej gazu dielektrycznego (tj. jego ciśnienia i kompozycji). Aby osiągnąć ten cel, musi zostać zbudowane proste urządzenie sterujące zdolne do błyskawicznych zmian parametrów tego gazu. Po dodaniu tego urządzenia do głównej konstrukcji komory, będzie ono w stanie efektywnie sterować częstotliwościami pulsowań jej pola.

5. Wyzwolenie zjawisk jakich zadaniem jest **odzyskanie ciepła** rozpraszanego przez iskry (w ten sposób wyeliminowanie strat energii następujących podczas działania komory). Celem tego etapu jest tak zmienić przebiegi procesów zachodzących w działającej komorze, aby spowodowały one zamianę energii cieplnej zawartej w gorącym gazie dielektrycznym w ładunki elektryczne gromadzące się na elektrodach komory. Aby osiągnąć ten cel całkowite zrozumienie złożonych procesów zachodzących w komorze musi zostać osiągnięte, zaś potem dokonane zostać musi przekształcenie tych zjawisk w wymaganym kierunku tak aby wynikowa komora czyniła użytek z możliwości efektu telekinetycznego (patrz opis tego efektu zawarty w podrozdziale H6.1, oraz opis jego wykorzystania z podrozdziału H6.2.1).

6. **Neutralizacja sił elektromagnetycznych** jakie działają na fizyczną konstrukcję (ścianki) komory. Celem tego etapu jest znaleźć taki wzajemny stosunek pomiędzy parametrami konstrukcyjnymi i parametrami pracy komory, że konstrukcja komory zostanie całkowicie uwolniona od akcji sił wytwarzanych podczas jej działania. Droga do osiągnięcia tego celu prowadzi poprzez stopniową zmianę parametrów konstrukcyjnych i operacyjnych komory oraz obserwowanie jaki wpływ wywierają te parametry na działanie sił występujących w komorze. Następnie konieczne będzie wybranie takich optymalnych wartości tych parametrów jakie spełnią cel etapu całkowicie uwalniając konstrukcję komory od działających w niej sił.

7. **Zbudowanie konfiguracji krzyżowej** lub nawet kapsuły dwukomorowej. Celem tego etapu jest takie zestawienie pojedynczych komór oscylacyjnych, aby razem pracowały

one jako konfiguracja krzyżowa lub nawet kapsuła dwukomorowa. Osiągnięcie tego celu wymaga dokonania różnorodnych zmian i dopasowań w sterowaniu komór składowych, jak również w zjawiskach w nich zachodzących, tak że wynikowa konfiguracja będzie pracowała efektywnie jako całość i pozostanie przy tym całkowicie sterowalna.

8. Nieograniczone zwiększanie zasobów energii komory. Celem tego etapu jest eksperymentalne wykrycie i usunięcie wszelkich możliwych przeszkód jakie mogłyby ograniczać ilość energii akumulowanej w zbudowanej poprzednio konfiguracji krzyżowej lub kapsule dwukomorowej. Docelowym poziomem upakowania energii w komorze jaki powinien zostać osiągnięty na tym etapie jest około dziesięciokrotne przekroczenie wartości strumienia startu przez rozpracowywaną konfigurację komór. Osiągnięcie tego celu będzie dosyć trudnym zadaniem, jako iż badania będą wymagały niezwyklej ostrożności i działań zabezpieczających, ponieważ przeładowane energią magnetyczną komory w razie uszkodzenia będą eksplodowały ze siłą potężnych bomb termojądrowych. Dla przykładu kapsuła dwukomorowa o objętości 1 metra sześciennego wypełniona polem magnetycznym o wartości dziesięciokrotnie przewyższającej jej strumień startu może eksplodować ze siłą około 10 megaton TNT. Wywołane przez nią zniszczenie byłoby więc równe prawie połowie zniszczenia od eksplozji tunguskiej na Syberii z 1908 roku, przez fachowców ocenianej na około 30 megaton TNT.

9. Modulowanie myślami ludzkimi pulsacji pola wytwarzanego przez komorę - patrz opisy z rozdziału N. W efekcie takiego modulowania komora uzyska dodatkową możliwość działania jako nadawczo-odbiorcza stacja telepatyczna (telepatyzer).

10. Zbudowanie komory drugiej generacji. Po opanowaniu budowy niezawodnych komór pierwszej generacji oraz dogłębnym poznaniu zjawisk i oczywistych niebezpieczeństw kryjących się w napełnianiu ich energią, a także po gruntownym przebadaniu wszystkich aspektów technicznego użycia telekinezy, możliwe będzie rozważenie podjęcia eksperymentów nad komorami drugiej generacji (tj. komorami samozapełniającymi się energią). Kierunek działania wskazywać będą baterie telekinetyczne - patrz podrozdział K2.4. Jednak zalecam tu szczególną ostrożność i **przestrzegam przed pochopnym podjęciem takich badań**, zanim niezawodne konstrukcje komór pierwszej generacji zostaną wypracowane. Wszakże każda taka samozapełniająca się energią komora, będzie także samo-uzbrajającą się bombą o niewyobrażalnej mocy zniszczenia.

Przeglądając powyższy program budowy komory oscylacyjnej zapewne narzuci się spostrzeżenie iż aż do końca etapu 3 celowo został on posegmentowany na szereg "małych kroków", w moim założeniu wystarczająco prostych aby stanowić wykonalne zadanie dla pojedynczego badacza. Stąd też program ten może zostać stopniowo realizowany zarówno przez indywidualnych hobbystów, jak i przez niewielkie zespoły rozwojowe. Szczególnie zaś nadaje się on do realizacji jako ciąg tematów dyplomowych dla studentów ostatniego roku uczelni lub szkół technicznych o profilu elektrycznym (lub elektronicznym). Dla przykładu etapy 1(a), 1(b), 1(c), 2(a), 3(a) już obecnie stanowią gotowe tematy prac dyplomowych wystarczająco prostych aby być skompletowanymi w przeciętnych laboratoriach uczelnianych lub przyszkolnych. (Trochę tu szkoda, że swoje badania zmuszony jestem wykonywać w całkowitej konspiracji, gdyby bowiem - jak inni naukowcy specjalizujący się w powszechnie uznanej już tematyce, posiadał swobodę otwartego prowadzenia swoich badań, wtedy ja sam mógłbym zrealizować powyższy program budowy. W takim zaś przypadku działająca komora oscylacyjna zapewne już od dawna służyłaby ludzkości. Niestety, aby przeżyć muszę oficjalnie wykonywać to co jest odemnie wymagane, zaś pozostający mi do dyspozycji czas prywatny w całości pochłaniany jest przez badania teoretyczne i publikowanie ich wyników. Na realizację więc w konspiracji również i budowy komory nie posiadam już ani warunków ani czasu.)

Na zakończenie tego podrozdziału warto tu podkreślić, iż po skompletowaniu etapu 3 prototypy komory oscylacyjnej zaczną być użyteczne przemysłowo, jako że z powodzeniem będą już wtedy mogły one wygrywać współzawodnictwo w różnorodnych zastosowaniach z ciężkimi i nieporęcznymi elektromagnesami. Dlatego też począwszy od etapu numer 4,

rozpracowywana komora oscylacyjna nabędzie zdolności do zarabiania na sobie i w ten sposób opłacania swojego dalszego rozwoju. Również począwszy od etapu 4 urządzenie to szybko rozprzestrzeni się na świecie i przejmie na siebie różnorodne funkcje jakie dotychczas wypełniane są przez inne urządzenia - patrz podrozdział C9.

C8.3. Przykłady tematów badawczych inicjujących prace nad komorą

Jak to wynika z podrozdziału C8.2 pierwsze trzy etapy budowy komory oscylacyjnej mogą z powodzeniem zostać zrealizowane nawet przez pojedynczego badacza. Z kolei po ich skompletowaniu komora zacznie przynosić dochód, sama więc zacznie finansować swój dalszy rozwój. Stąd przy odrobinie szczęścia i talentu wynalazczego, osoba jaka obecnie zainwestuje w owo urządzenie, być może już wkrótce posiadać klucz do całej energii naszej planety. Jest to niewypowiedzianie duża stawka do wygrania, zaś rodzaj wkładu wymagany na początku aby włączyć się do gry dostępny jest praktycznie dla każdego. Każdy bowiem może gdzieś zdobyć kilka płytek pleksi, paczkę szpilek krawieckich, jakieś kondensatory i cewki, starą maszynę Wimshursta albo cewkę zapłonową z akumulatorkiem. Co na obecnym etapie jest najbardziej potrzebne to dedykacja, dużo zdrowego rozsądku, smykałka wynalazcza, oraz sławna w świecie zdolność Polaków do efektywnej improwizacji. Dlaczegoż więc nie spróbować?

W celu ułatwienia takiego startu, w niniejszym podrozdziale przytoczyłem kilka początkowych eksperymentów nad komorą oscylacyjną. Mogą one być użyte jako tematy badawcze realizowane samodzielnie przez indywidulanych hobbystów, albo też jako tematy prostych prac dyplomowych wydawanych studentom różnych uczelni lub szkół średnich. Oto one:

Temat 1: "Eksperymentalne badania bezcewkowych obwodów oscylacyjnych". Dotychczasowe obwody oscylacyjne zawsze zaopatrywane były w cewkę dostarczającą im wymaganej inercji elektrycznej (induktancji). Jednakże wykorzystywany w działaniu cewek przepływ prądu przez zwoje przewodnika nie jest jedynym znanym zjawiskiem zdolnym do dostarczenia wymaganej inercji elektrycznej. Innym dobrze znanym takim zjawiskiem jest zwykła iskra elektryczna. Stąd istnieje możliwość że odpowiednio zaprojektowane snopy wielu iskier elektrycznych przeskakujących po równoległych trajektoriach w niektórych obwodach oscylacyjnych będą w stanie zastąpić cewki indukcyjne. Obwód jaki najlepiej spełniałby wszystkie warunki nakładane przy takim zastępowaniu to konwencjonalny obwód Henry'ego. Jego cechą jest bowiem, że dla prawidłowego działania wymaga on obecności elektrod produkujących iskry. Stąd iskry stanowią w nim naturalną manifestację dostarczanej przez niego oscylacyjnej odpowiedzi na początkowe naładowanie elektryczne.

Celem niniejszego tematu badawczego jest takie zmodyfikowanie tradycyjnego obwodu Henry'ego aby zdolny był on do dostarczenia oscylacyjnej odpowiedzi wyłącznie w efekcie inercji snopów wytwarzanych przez siebie iskier elektrycznych i całkowicie bez użycia zewnętrznej cewki roboczej.

Badania te są eksperymentalne i obejmują one (1) zbudowanie obwodu oscylacyjnego Henry'ego, (2) skompletowanie na nim wymaganych badań, (3) zmodyfikowanie tego obwodu i powtarzanie badań aż nałożony cel końcowy zostanie osiągnięty.

Eksperymenty powinny być prowadzone na obwodzie oscylacyjnym Henry'ego (zbudowanym przez osobę podejmującą ten temat) który jest relatywnie prosty do zbudowania i modyfikowania.

W przypadku sukcesu w zrealizowaniu celu tych badań, osiągnięte wyniki posiadałyby istotne znaczenie naukowe i mogłyby dostarczyć danych dla przygotowania publikacji naukowej.

Cele, sposoby ich osiągnięcia, oraz podłoże naukowe dla niniejszego tematu badawczego opisane zostały w podrozdziale C8.2 (patrz etap 1 (c) opisanych tam eksperymentów).

Temat 2: "System do samoregulacji 90 przesunięcia fazowego w układzie dwóch obwodów oscylacyjnych z iskrownikiem". Jak dotychczas pole magnetyczne wytwarzane jest na Ziemi przy użyciu tylko jednej zasady działania zrealizowanej w postaci urządzenia zwanego elektromagnesem. Zasada ta posiada jednakże wiele ograniczeń i wad wrodzonych które powodują iż uzyskiwane dotychczas wydatki pola są stosunkowo niskie i niewystarczające dla wielu istotnych zastosowań praktycznych (np. dla napędu wehikułów latających). Z tego też powodu ostatnio podjęte zostały prace badawcze nad zupełnie inną zasadą wytwarzania pola, która zrealizowana zostanie w urządzeniu zwanym "komora oscylacyjna". W zasadzie tej źródłem pola będzie iskra elektryczna rotująca po obwodzie kwadratu. Jednym z problemów czekających rozwiązania już w początkowej fazie realizacji tej zasady jest takie samoregulowanie dwóch niezależnych od siebie obwodów oscylacyjnych z iskrownikiem, aby wytwarzane przez nie iskry przeskakiwały ze wzajemnym przesunięciem fazowym wynoszącym 90 stopni. Niniejszy temat badawczy służy właśnie próbie eksperymentalnego wypracowania takiego prostego systemu samoregulującego.

Celem tego tematu jest takie zmodyfikowanie dwóch konwencjonalnych obwodów oscylacyjnych Henry'ego aby zdolne one były do samoregulacji 90 przesunięcia fazowego w wytwarzanych przez nie iskrach (przeskakujących po dwóch prostopadłych i nawzajem krzyżujących się trajektoriach).

Badania te są eksperymentalne i obejmuje one (1) zbudowanie dwóch obwodów oscylacyjnych Henry'ego tak aby formowały one "model komory oscylacyjnej" opisany w punkcie 2 (a) rozdziału C8.2 tej monografii, (2) skompletowanie na nich wymaganych badań, (3) zmodyfikowanie tych obwodów i powtarzanie badań aż nałożony cel końcowy zostanie osiągnięty.

Eksperymenty powinny być prowadzone na dwóch nawzajem skrzyżowanych obwodach oscylacyjnych Henry'ego opisanych dla etapu 1 (a) z rozdziału C8.2 tej monografii (zbudowanych przez osobę podejmującą ten temat) które są relatywnie proste do zbudowania i modyfikowania.

W przypadku sukcesu z zrealizowaniem celu tych badań, osiągnięte wyniki posiadałyby istotne znaczenie naukowe i mogłyby dostarczyć danych dla przygotowania publikacji naukowej.

Cele, sposoby ich osiągnięcia, oraz podłoże naukowe dla niniejszego tematu badawczego opisane zostały w podrozdziale C8.2 (patrz etap 2 (a) opisanych tam eksperymentów).

Temat 3: "Magnetyczne zasilanie obwodów oscylacyjnych z iskrownikiem". Większość dotychczasowych obwodów oscylacyjnych zasilana jest w energię za pomocą impulsów elektrycznych. Jednakże w niektórych przypadkach znacznie korzystniejsze byłoby ich zasilanie impulsami pola magnetycznego dostarczanego w sposób bezdotykowy. Dla przykładu takie zasilanie za pośrednictwem sprzężenia magnetycznego umożliwiałoby wymianę energii pomiędzy obwodami drgającymi o różniących się parametrach pracy (np. różniących się napięciem czy częstotliwościami drgań).

Celem tego tematu jest zbudowanie możliwie najprostszego zasilacza magnetycznego, jaki za pomocą impulsów pola bezdotykowo zasilaby w energię tradycyjny obwód oscylacyjny Henry'ego.

Temat ten jest eksperymentalny i obejmuje on (1) zbudowanie obwodu oscylacyjnego Henry'ego stanowiącego obiekt zasilania w energię, (2) zbudowanie możliwie najprostszego zasilacza jaki bezdotykowo dostarczałby energię do obwodu Henry'ego za pomocą impulsów magnetycznych, (3) badania obu urządzeń mające na celu ustalenie warunków efektywnego przepływu energii od zasilacza do obwodu Henry'ego, (4) takie modyfikowanie i usprawnianie zasilacza oraz powtarzanie badań aż wynikowe urządzenie zasilające będzie w stanie samoczynnie wzbudzić oscylacje iskrowe w obwodzie Henry'ego.

W raporcie końcowym z tych badań, oprócz jej przebiegu i efektów, uwypuklone też powinny zostać znalezione warunki najefektywniejszego przepływu energii z zasilacza do obwodu Henry'ego, oraz praktycznie użyte sposoby spełnienia tych warunków. Dzięki temu

badania te oraz zasilacz zbudowany w ich efekcie stanowiłyby eksperyment początkowy (pilotujący) dla bardziej zaawansowanych badań prowadzonych w etapie (roku) następnym. Te dalsze badania ukierunkowane byłyby na podniesienie efektywności, wydajności i uniwersalności (np. zakresu częstotliwości roboczych) opracowywanego urządzenia zasilającego.

W przypadku sukcesu w zrealizowaniu celu tych badań, osiągnięte wyniki posiadałyby istotne znaczenie naukowe i mogłyby dostarczyć danych dla przygotowania publikacji naukowej.

Cele, sposoby ich osiągnięcia, oraz podłoże naukowe dla niniejszego tematu badawczego opisane zostały w podrozdziale C8.2 (patrz etap 3 (a) opisanych tam eksperymentów).

C9. Przyszłe zastosowania komory oscylacyjnej

Ponieważ komora oscylacyjna jest aż tak zaawansowanym i aż tak uniwersalnym akumulatorem energii, po zbudowaniu będzie ona posiadała niezliczone zastosowania praktyczne. Najbardziej istotne z owych zastosowań są opisane w następnym rozdziale CB.

C10. Moje monografie poświęcone komorze oscylacyjnej

Zanim niniejsza praca została opublikowana, komora oscylacyjna była już prezentowana w kilku innych moich monografiach. Poniższy wykaz zestawia w porządku chronologicznym kolejność ukazywania się najważniejszych z nich (opublikowanych zostało znacznie więcej). Warto tu zaznaczyć, że monografie [5C] i [6C] posiadają też swoje polskojęzyczne odpowiedniki (patrz monografie oznaczone [5] i [6] w rozdziale Y, oraz [1P2.2] i [2P2.2] z podrozdziału P2.2).

[1C] "Theory of the Magnocraft". Zawierała ona pierwszy krótki opis komory oscylacyjnej (jeden niewielki rozdział). Publikowana ona była w języku angielskim w następujących wydaniach:

(a) Pierwsze wydanie nowozelandzkie, styczeń 1984 roku, ISBN 0-9597698-0-3.

(b) Pierwsze wydanie USA, czerwiec 1985 roku - opublikowane w USA przez: Energy Unlimited, P.O. Box 35637 Sta. D, Albuquerque, NM 78176.

(c) Pierwsze wydanie polskie (tj. napisane w języku polskim - patrz [1]), zatytułowane "Teoria Magnokraftu", Invercargill, Nowa Zelandia, marzec 1986, ISBN 0-9597698-5-4; 136 stron, 58 rysunków. (Jak to wyjaśniłem w podrozdziale A4, to właśnie ową monografię [1] "Teoria Magnokraftu", Rada Naukowa Instytutu Technologii Budowy Maszyn Politechniki Wrocławskiej w 1986 roku zdyskwalifikowała jako nie kwalifikującą się na propozycję zakresu tematycznego dla mojej rozprawy habilitacyjnej.)

(d) Drugie wydanie nowozelandzkie - poszerzone, Invercargill, sierpień 1984 roku, ISBN 0-9597698-1-1; 110 stron plus 53 ilustracji.

[2C] "The Oscillatory Chamber - a breakthrough in the principles of magnetic field production". Była to pierwsza angielskojęzyczna monografia w całości poświęcona opisowi komory oscylacyjnej. Opublikowana ona została w następujących wydaniach:

(a) Pierwsze wydanie nowozelandzkie, grudzień 1984 roku, ISBN 0-9597698-2-X.

(b) Pierwsze wydanie USA, opublikowane w magazynie "Energy Unlimited", Issue 19/1985, strony 15 do 43. To specjalne wydanie magazynu (opublikowane przez "Energy Unlimited", P.O. Box 35637, Station D, Albuquerque, NM 87176, USA) przedrukowało całą monografię o komorze oscylacyjnej.

(c) Pierwsze wydanie zachodnio-niemieckie (w języku niemieckim) zatytułowane, "Die 'Schwingkammer' Energie & Antrieb fur das Weltraumzeitalter", opublikowane przez: Raum & Zeit Verlag, Dammtor 6, D-3007 Gehrden, West Germany; czerwiec 1985 roku, ISBN 3-89005-006-9; 64 strony (włączając w to 7 rysunków).

(d) Drugie wydanie nowozelandzkie, przerobione, Invercargill, październik 1985 roku, ISBN 0-9597698-4-6; 115 stron plus 15 ilustracji. Wydanie to zawierało także pierwszą prezentację Konceptu Dipolarnej Grawitacji.

[3C] "The Magnocraft: a saucer-shaped space vehicle propelled by a pulsating magnetic field", Invercargill, Nowa Zelandia, 1986 rok, ISBN 0-9597698-3-8; 300 stron. Stanowiła ona poszerzenie, udoskonalenie i uaktualizowanie monografii [1C]. Cały jeden jej rozdział poświęcony był komorze oscylacyjnej.

[4C] "The Magnocraft - Earth's version of a UFO", Treatise, Dunedin, New Zealand 1990, ISBN 0-9597698-6-2, 420 stron (włączając w to 7 tablic i 163 ilustracji). Przedpublikacyjne wersje tej monografii upowszechniane były już od 1987 roku. Była ona poszerzeniem, udoskonaleniem i uaktualizowaniem monografii [3C].

[5C] "Tapanui Cataclysm - an explanation for the mysterious explosion in Otago, New Zealand, 1178 A.D.". Dunedin, New Zealand, 1989 rok, ISBN 0-9597698-7-0, a private edition by the author (62 strony włączając 26 ilustracji). Monografia ta posiadała potem kilka dalszych wydań - patrz [2O4.2] i [5].

[6C] "The magnetic extraction of energy from the environment". Dunedin, New Zealand, 1990 rok, ISBN 0-9597946-1-1; 38 stron (włączając 14 ilustracji).

[7C] "Advanced Magnetic Propulsion Systems", Treatise, Dunedin, New Zealand, październik 1990 roku, ISBN 0-9597698-9-7, 460 stron (włączając 7 tablic i 163 ilustracji). Jest to dalsze udoskonalenie monografii [4C] oraz angielskojęzyczny poprzednik niniejszej monografii.

[8C] "Komora oscylacyjna czyli magnes jaki wzniesie nas do gwiazd", Dunedin 1994 roku, ISBN 0-9597946-2-X, 184 stron (w tym 4 tablice i 39 rysunków). Jest to uaktualizowane wydanie polskojęzyczne monografii [2C].

[9C] "The Oscillatory Chamber, arkway to the stars", Dunedin, New Zealand, 1994 rok, ISBN 0-9583380-0-0, 365 stron tekstu plus 104 ilustracje i 7 tablic. Jest to trzecie angielskojęzyczne wydanie monografii [2C].

Monografie [5C] i [6C] podsumowują opisy komory oscylacyjnej, nie zawierają jednak dokładniejszego wyjaśnienia dla jej konstrukcji i zasady działania.

C11. Symbole, notacje i jednostki występujące w rozdziale C

symbole – wyjaśnienia	[jednostki]
a - wymiar boczny sześcianu	[metr]
A - pole powierzchni	[metr·metr]
Č - siła ściskająca	[Newton]
C - pojemność elektryczna	[Farad]
E - symbol oznaczający "elektroda"	[-]
f - częstotliwość pulsowania	[1/second]
F - strumień magnetyczny	[Weber]
F _o - składowa stała strumienia magnetycznego	[Weber]
i - natężenie prądu elektrycznego	[Ampere]
l - odległość lub długość	[meter]
L – indukcyjność	[Henry]
m - symbol oznaczający "oś magnetyczna"	[-]
M - siła magnetyczna oddziaływująca na prąd elektr.	[Newton]
n - liczba zwojów cewki w jednostce jej długości	[-]
p - liczba segmentów wydzielonych w elektrodzie	[-]
P - segment elektrody	[-]
q - ładunek elektryczny	[Coulomb]
R - oporność elektryczna przewodzenia	[Ohm]
s - współczynnik mobilności iskry	[-]
S - symbol oznaczający "iskra"	[-]
t – czas	[second]
T - okres pulsowań	[second]
T̄ - siła rozciągająca	[Newton]
U - napięcie zapoczątkowujące wyładowanie w komorze	[Volt]
ΔF- amplituda pulsowań strumienia magnetycznego	[Weber]
ε - stała dielektryczna dla gazu wypełniającego komorę	[Farad/metr]
μ - przenikalność magnetyczna gazu dielektrycznego	[Henry/metr]
Ω - oporność elektryczna gazu dielektrycznego w komorze wyznaczona w momencie początku przeskoku iskry	[Ohm*metr]

Indeksy przyporządkowane elektrodom:

- B - Odnosi się do tylnej elektrody (back)
- F - Odnosi się do przedniej elektrody (front)
- L - Odnosi się do lewej elektrody (left)
- R - Odnosi się do prawej elektrody (right)

Indeksy przyporządkowane komorom oscylacyjnym:

- I - Odnosi się do wewnętrznej (inner) komory oscylacyjnej
- O - Odnosi się do zewnętrznej (outer) komory oscylacyjnej

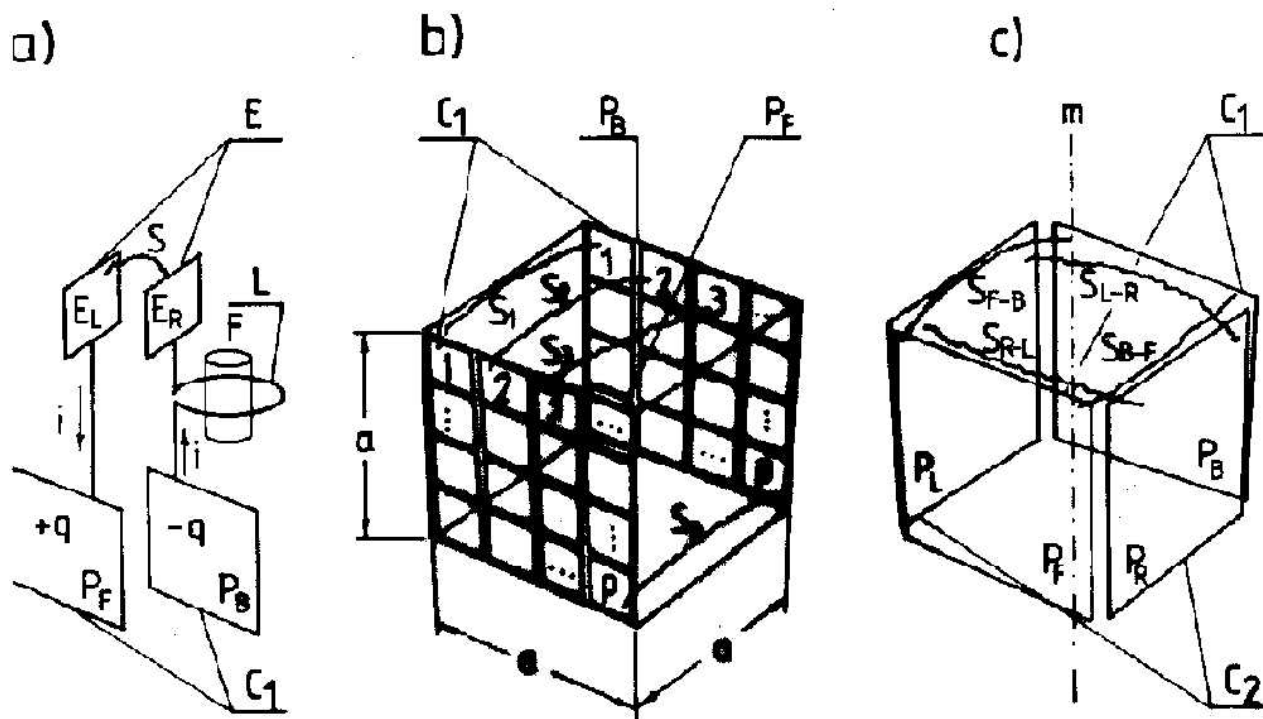
Indeksy przyporządkowane wielkościom fizycznym:

- N - Odnosi się do północnego bieguna magnetycznego (north)
- S - Odnosi się do południowego bieguna magnetycznego (south)
- C - Oznacza strumień krążący (circulating flux) kapsuły dwukomorowej
- R - Oznacza strumień wynikowy (resultant flux) kapsuły dwukomorowej

C-76

Tablica C1. Wykorzystanie komory oscylacyjnej. Zestawiono tu kilka przykładów obecnych urządzeń energetycznych, które w przyszłości zastąpione zostaną przez komorę oscylacyjną z uwagi na jej zdolność do wielo-wymiarowej transformacji energii. (Zauważ, że przykłady wielu dalszych urządzeń jakie zapewne też zastąpione kiedyś będą przez komory oscylacyjne omówiono w podrozdziale C9.)

Nr	Urządzenie zastępowane przez komorę	Rodzaj energii		zastępujących dany rodzaj urządzenia
		Dostarczonej	Uzyskanej	
1.	Elektromagnes	Prąd elektryczny	Pole magnetyczne	Energia elektryczna dostarczona do komory przetransformowana zostaje na pole magnetyczne.
2.	Grzejnik elektryczny	Prąd elektryczny	Ciepło	Gorący gaz dielektryczny z komory zostaje przepompowywany poprzez wymiennik ciepła.
3.	Silnik elektryczny	Prąd elektryczny	Ruch mechaniczny	Fale sterowanego pola magnetycznego wytwarzanego przez układ komór oscylacyjnych stojana powodują ruch mechaniczny przewodzącego wirnika.
4.	Transformator	Prąd elektryczny	Prąd elektryczny o innych parametrach	Dwie komory o różniących się parametrach pracy wymieniają energię za pośrednictwem pola magnetycznego (tj. sterując przesunięciem fazowym wytwarzanego przez siebie pola).
5.	Silnik spalinowy	Ciepło	Ruch mechaniczny	Podgrzewanie gazu dielektrycznego dostarcza komorze energii wykorzystywanej do produkcji ruchu mechanicznego jak w silniku elektrycznym.
6.	Ogniwo termoelektryczne	Ciepło	Prąd elektryczny	Podgrzewanie gazu dielektrycznego zwiększa energię komory. Energia ta, przetransformowana przez komorę na ładunki elektryczne, może być z niej podjęta w postaci prądu elektrycznego.
7.	Generator elektryczności	Ruch mechaniczny	Prąd elektryczny	Przemieszczanie jednej komory w zasięgu pola innej wytwarza energię oddziaływania pól jaka potem może zostać podjęta w postaci prądu.

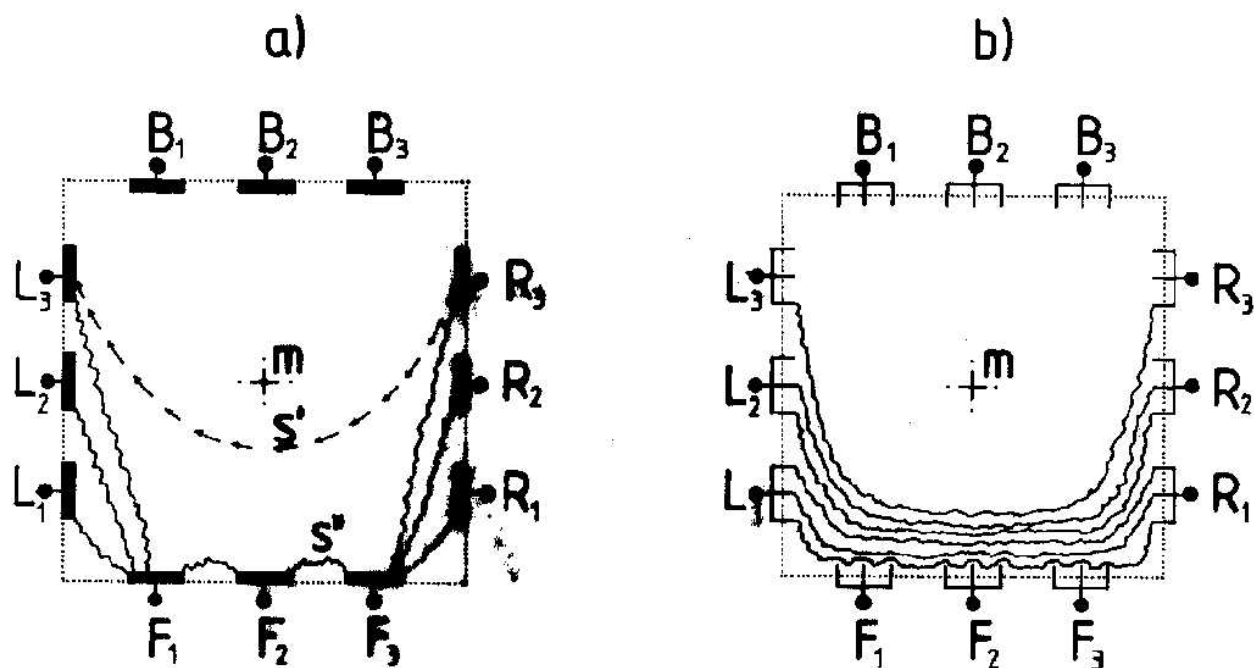


Rys. C1. Formowanie komory oscylacyjnej. Trzy części tego rysunku pokazują trzy kolejne etapy przekształcania konwencjonalnego obwodu oscylacyjnego z iskrownikiem w komorę oscylacyjną.

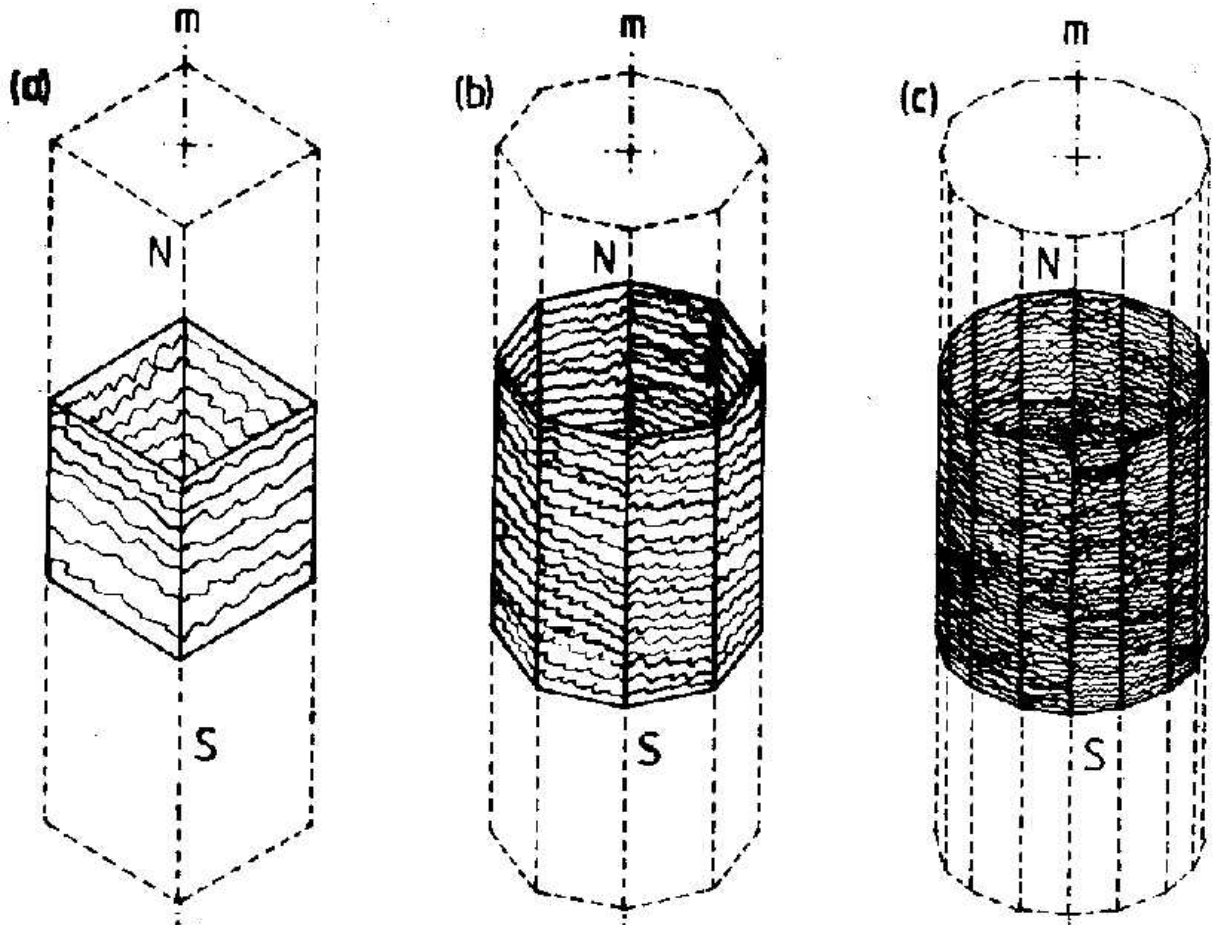
(a) Tradycyjna postać obwodu oscylacyjnego z iskrownikiem, jaka wynaleziona została przez J. Henry'ego w 1845 roku. Trzy istotne składowe tego obwodu (tj. pojemność " C_1 ", indukcyjność " L " i przerwa iskrowa " E ") dostarczane są przez trzy oddzielne urządzenia.

(b) Zmodyfikowana wersja obwodu oscylacyjnego " C_1 " z iskrownikiem. Wszystkie trzy jego istotne składowe zostały tu skoncentrowane w jednym urządzeniu, tj. układzie dwóch przewodzących elektrod " P_F " i " P_B " przymocowanych do dwóch przeciwległych ścianek komory sześcienniej wykonanej z materiału izolacyjnego. Obie elektrody " P_F " i " P_B " podzielone z kolei zostały na kilka pooddzielanych od siebie segmentów oznaczonych numerami "1, 2, ..., p". Długość boku wynikowej komory sześcienniej z owymi elektrodami w środku oznaczona została przez " a ".

(c) Komora oscylacyjna uformowana poprzez zestawienie razem dwóch zmodyfikowanych obwodów " C_1 " i " C_2 " identycznych do obwodu pokazanego w części (b) tego rysunku. Kolejne pojawianie się pęków iskier oznaczonych przez " S_{R-L} ", " S_{F-B} ", " S_{L-R} ", " S_{B-F} " jakie zawsze przeskakują wzdłuż powierzchni bocznych ścianek komory leżących po ich lewych stronach, powoduje wytworzenie rodzaju łuku elektrycznego rotującego naokoło obwodu komory i wytwarzającego potężne pole magnetyczne.



Rys. C2. Uzasadnienie użycia igło-kształtnych elektrod. Rysunek pokazuje odgórny widok dwóch wersji komory oscylacyjnej podczas ich działania. W obu wersjach pęki iskier pokazane zostały w procesie przesłakiwania z elektrod oznaczonych jako "R" (tj. prawych - "right") do elektrod oznaczonych jako "L" (lewych - "left"). Ponieważ we wnętrzu komory, wzdłuż jej pionowej osi "m" panuje silne pole magnetyczne, skaczące iskry zostają przyparte ku powierzchni elektrody oznaczonej jako "F" (tj. przedniej - "front"). To przypieranie powoduje, iż lewa komora (a) wykorzystująca płyto-kształtne elektrody, zamiast pożądanego przebiegu (s') swoich iskier, uzyskuje ten przebieg wzdłuż linii najmniejszego oporu (s'') prowadzącego przez materiał przednich płyt "F₁", "F₂", "F₃". Jednakże takie przeskoki "na skróty" nie są możliwe w komorze prawej (b) z igło-kształtnymi elektrodami, gdzie ostre końce elektrod igłowych odpychają iskry czyniąc niemożliwym ich wnikanie do materiału elektrod "F".

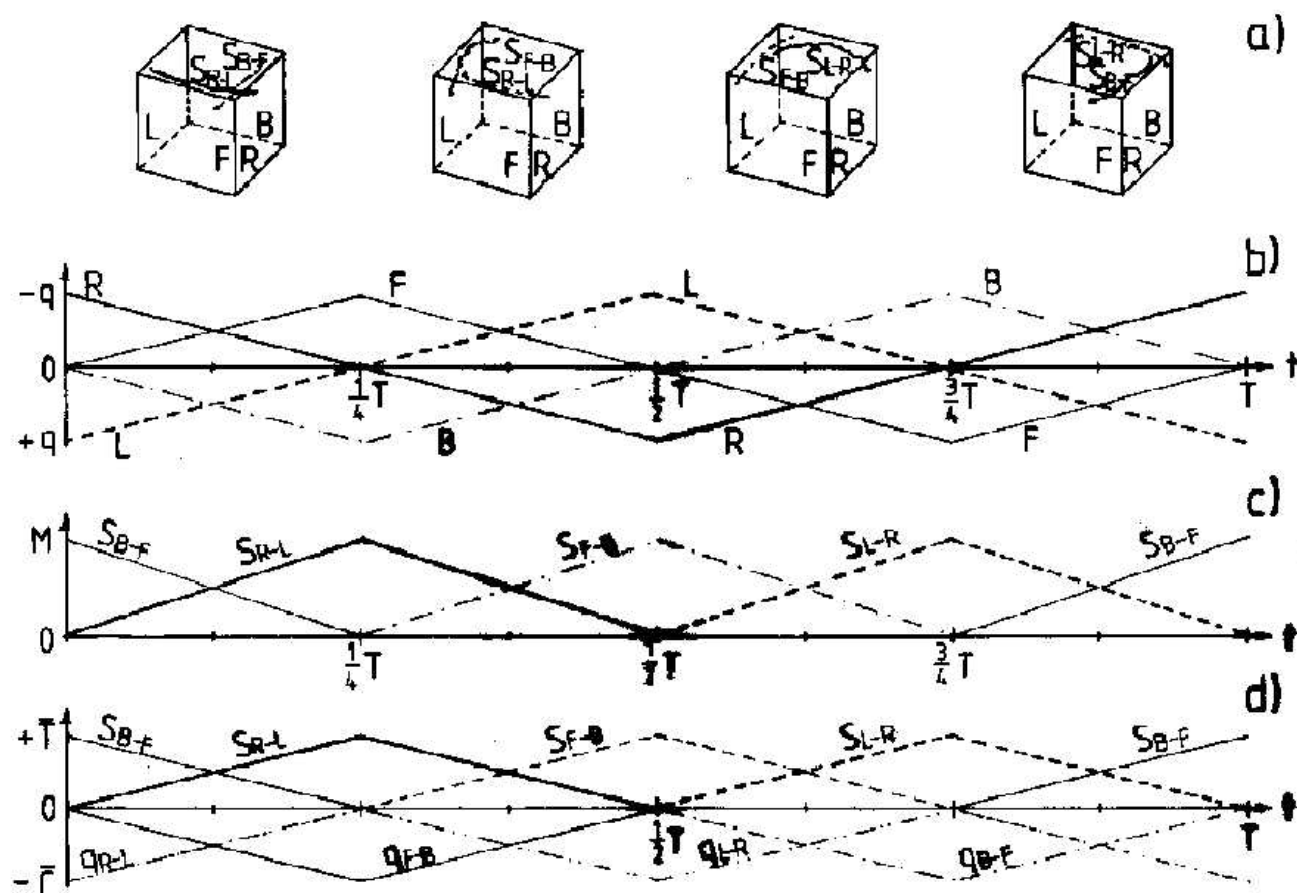


Rys. C3. Przewidywany wygląd komór oscylacyjnych (a) pierwszej, (b) drugiej i (c) trzeciej generacji. Zauważ że powyższy rysunek pokazuje wygląd pojedynczych komór, podczas gdy wygląd nieco podobnych do nich kasul dwukomorowych pokazany został na rysunkach C5, C6 i C8. Komory oscylacyjne przyjmowały będą niepozorną postać szklanej kostki lub kryształowego pręta. Ich przekrój poprzeczny, a tym samym i kształt zewnętrzny, zależał będzie od generacji do której dana komora należy. Komory pierwszej generacji (a) będą posiadały przekrój kwadratowy, drugiej generacji (b) - ośmioboczny, zaś trzeciej generacji (c) - szesnastoboczny. Pasma jasnych, zygzakowatych, migotliwych iskier koloru złocistego będą przebiegały poziomo wokół obwodu wewnętrznego ich ścianek bocznych. Iskry te będą jakby zamrożone w tych samych pozycjach, aczkolwiek od czasu do czasu będą one raptownie przemieszczały swój przebieg jak kłębowisko węży owiniętych wokół swojej zdobyczy. Stąd działająca komora oscylacyjna będzie sprawiała wrażenie kryształu upakowanego żywą energią, lub nawet jakiejś istoty zajętej tajemniczymi czynnościami żywymi. W zależności od generacji komory, zygzakujące w niej iskry będą wyglądały nieco inaczej. W komorach pierwszej generacji iskry będą nierównej średnicy - jedne bardzo grube inne zaś bardzo cienkie, ukierunkowane dosyć chaotycznie, czasami rozdwarzające się, nieprecyzyjnie uwarstwione, oraz wyglądające jakby wykonane ręcznie przez kiepskiego kowala. W komorach drugiej generacji będą już jednakowej grubości jakby wykonane przez zegarmistrza na dokładnej maszynie, aczkolwiek ciągle czasami nachodzące na siebie. Natomiast w komorach trzeciej generacji będą one równiusieńkiej grubości, uporządkowane jedna obok drugiej, oraz wyglądające jakby ktoś precyzyjnie nawinał je naokoło wnętrza komory. Przerwane linie zaznaczają przebiegi kolumn pola magnetycznego wytwarzanego przez te komory. Kolumny te rozciągają się wzdłuż osi magnetycznych "m" komór. Jeśli komory oglądane będą z kierunku prostopadłego do linii sił tego pola (tj. dokładnie jak zostały one przedstawione na tym rysunku) wtedy owe kolumny pola będą przechwytywały światło i stąd powinny być widziane przez nieuzbrojone oko jako tzw. "czarne belki" rozprzestrzeniające się z komór w obu kierunkach i odzwierciedlające ich przekrój poprzeczny - patrz opisy takich belek przedstawione w podrozdziale F10.4. Pole to powinno także czynić wnętrze komór nieprzezroczystym. Stąd oglądane z boku powinny one wyglądać jakby wypełnione czarnym dymem. Gdy jednak oglądane są wzdłuż linii sił pola magnetycznego, prześwit przez komory powinien być przeźroczysty - za wyjątkiem przypadków pokazanych na rysunkach C6 i C8 kiedy dana komora zagina strumień krążący kapsuły dwukomorowej.

(a) Wygląd boczny komory oscylacyjnej pierwszej generacji w kształcie sześciangu - patrz też rysunek S6.

(b) Wygląd boczny komory drugiej generacji. Jej czoło ma kształt ośmioboku - patrz też rysunki P19(D) i P29.

(c) Wygląd boczny komory oscylacyjnej trzeciej generacji. Przy pierwszym wejrzeniu sprawia ona wrażenie kawałka niemal okrągłego pręta z kryształu, mieniaącego się w środku złotymi iskrami.



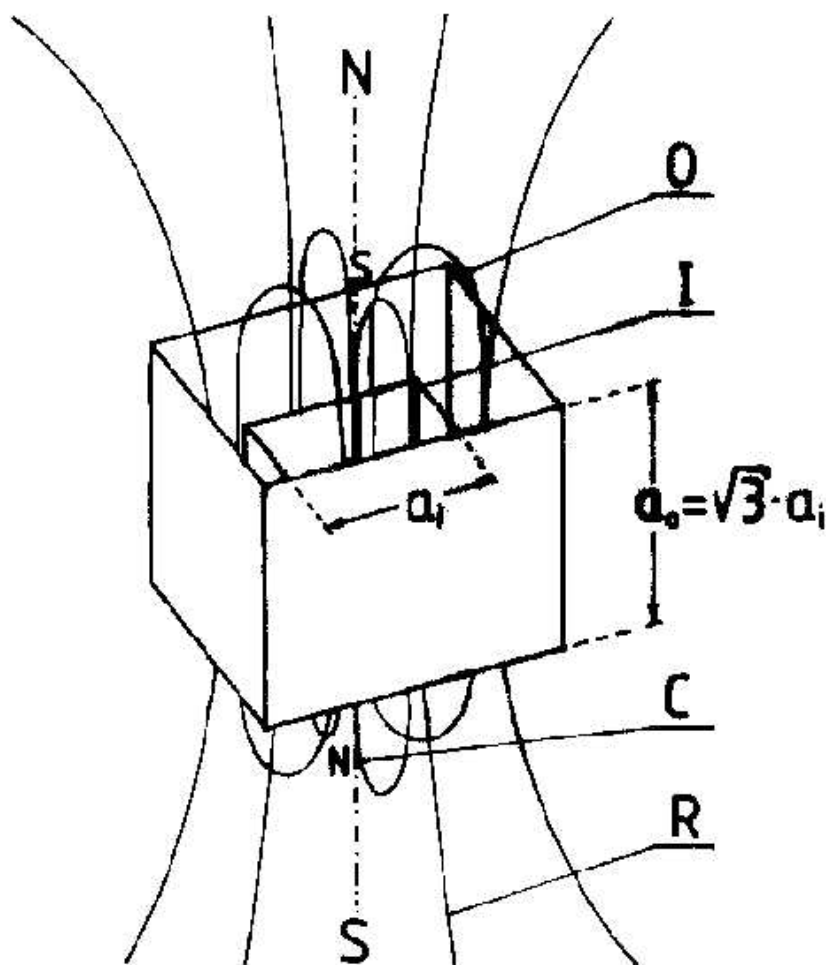
Rys. C4. Neutralizacja sił elektro-magnetycznych w komorze oscylacyjnej. Jej mechanizm wykorzystuje przeciwieństwo kierunków działania dwóch sił, tj. Coulomb'owskiego ściskania i elektromagnetycznego rozrywania, aby spowodować ich wzajemne zneutralizowanie się.

(a) Cztery główne stadia działania komory oscylacyjnej. Symbole: R, L, F, B - prawa, lewa, przednia i tylna elektrody komory (tj. right, left, front, back) jakie razem formują dwa współdziałające z sobą obwody oscylacyjne; S_{R-L} , S_{F-B} , S_{L-R} , S_{B-F} - cztery pęki iskier elektrycznych jakie pojawiają się kolejno po sobie podczas pojedynczego cyklu jej działania, formując w ten sposób jeden całkowity obieg łuku elektrycznego po kwadracie (iskry aktywne zaznaczono linią ciągłą, zaś iskry inercyjne linią przerywaną).

(b) Zmiany w potencjale elektrod podczas pełnego cyklu działania komory. Symbole: T - okres pulsowań komory; t - czas; +q, -q - dodatnie i ujemne ładunki elektryczne akumulujące się na elektrodach. Wzajemne przyciąganie się odwrotnych ładunków zgromadzonych na przeciwstawnych ściankach wywołują siły Coulomb'a ściskające komorę ku wewnątrz.

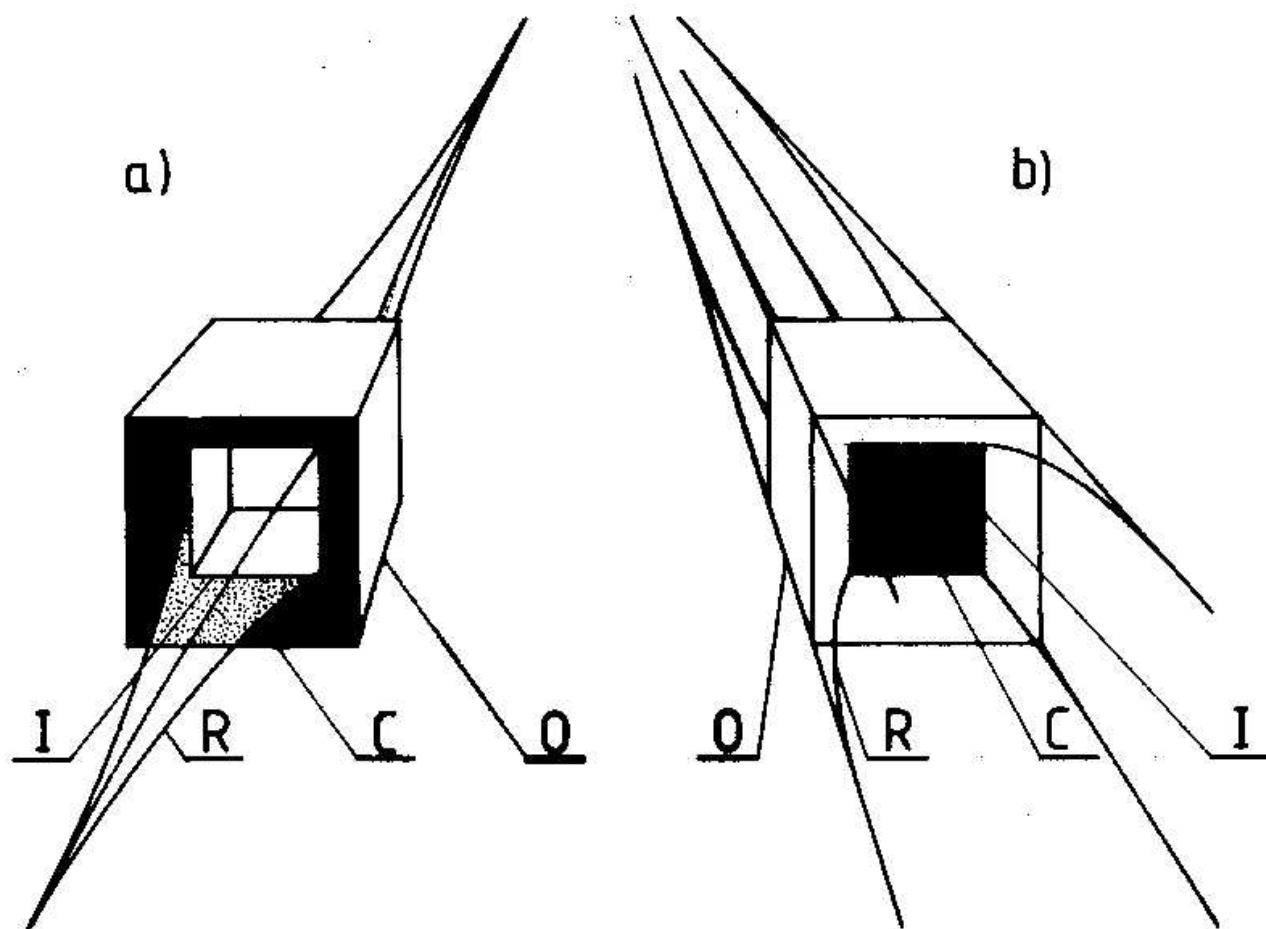
(c) Zmiany we wartości elektromagnetycznych sił odchylających (M) działających na kolejne iskry elektryczne. Siły te starają się rozerwać komorę ku zewnątrz, podobnie jak to czynią z elektromagnesami.

(d) Zmiany w siłach rozrywających "T" (tensing) komorę oraz siłach ściskających "C" (compressing) jakich działanie neutralizuje się nawzajem. Siły rozrywające "T" produkowane są przez elektromagnetyczne oddziaływania odchylające zachodzące pomiędzy iskrami i polem magnetycznym wypełniającym komorę. Siły ściskające "C" wywoływane są wzajemnym Coulomb'owskim przyciąganiem się przeciwstawnych ładunków elektrycznych zakumulowanych na przeciwległych elektrodach komory. Zauważ, iż obie te grupy sił muszą przyjmować symetryczny przebieg, ale przeciwstawne wartości. Stąd są one w stanie nawzajem znieść swoje działanie.



Rys. C5. Kapsuła dwukomorowa uformowana z komór oscylacyjnych pierwszej generacji. Jest to podstawowa konfiguracja dwóch komór oscylacyjnych, formowana w celu zwiększenia ich sterowalności. Powstaje ona poprzez osadzenie dwóch przeciwstawnie zorientowanych komór oscylacyjnych pierwszej generacji, jedna we wnętrzu drugiej. Z uwagi na potrzebę swobodnego "pływania" komory wewnętrznej (I) zawieszanej w środku komory zewnętrznej (O), boki "a" obu tych komór muszą wypełniać równanie (C9): $a_o = a_i \sqrt{3}$. Z powodu przeciwnego zorientowania biegunów magnetycznych obu komór kapsuły, wynikowe pole magnetyczne (R) odprowadzane z tej konfiguracji do otoczenia, stanowi algebraiczną różnicę pomiędzy wydatkami jej komór składowych. Zasada formowania takiego strumienia wynikowego została zilustrowana na rysunku C7. Kapsuły dwukomorowe umożliwiają łatwe sterowanie wszystkimi atrybutami wytwarzanego przez nie pola. Przedmiotem tego sterowania są następujące własności strumienia wynikowego (R): (1) moc pola - regulowana płynnie od zera do maksimum; (2) okres pulsowań (T) lub częstość pulsowań (f); (3) stosunek amplitudy pulsowań pola do jego składowej stałej ($\Delta F/F_o$ - patrz rysunek C12); (4) charakter pola, tj. czy jest ono stałe, pulsujące, czy przemienne; (5) krzywa zmian w czasie $F=f(t)$, np. czy jest to pole liniowe, sinusoidalne, czy zmieniane według "krzywej dudnienia"; (6) biegunowość (tj. z której strony kapsuły panuje biegun N lub biegun S).

Symbole: O - komora zewnętrzna (outer), I - komora wewnętrzna (inner), C - strumień krążący (circulating flux) uwięziony we wnętrzu kapsuły, R - strumień wynikowy (resultant flux) odprowadzany z kapsuły do otoczenia.

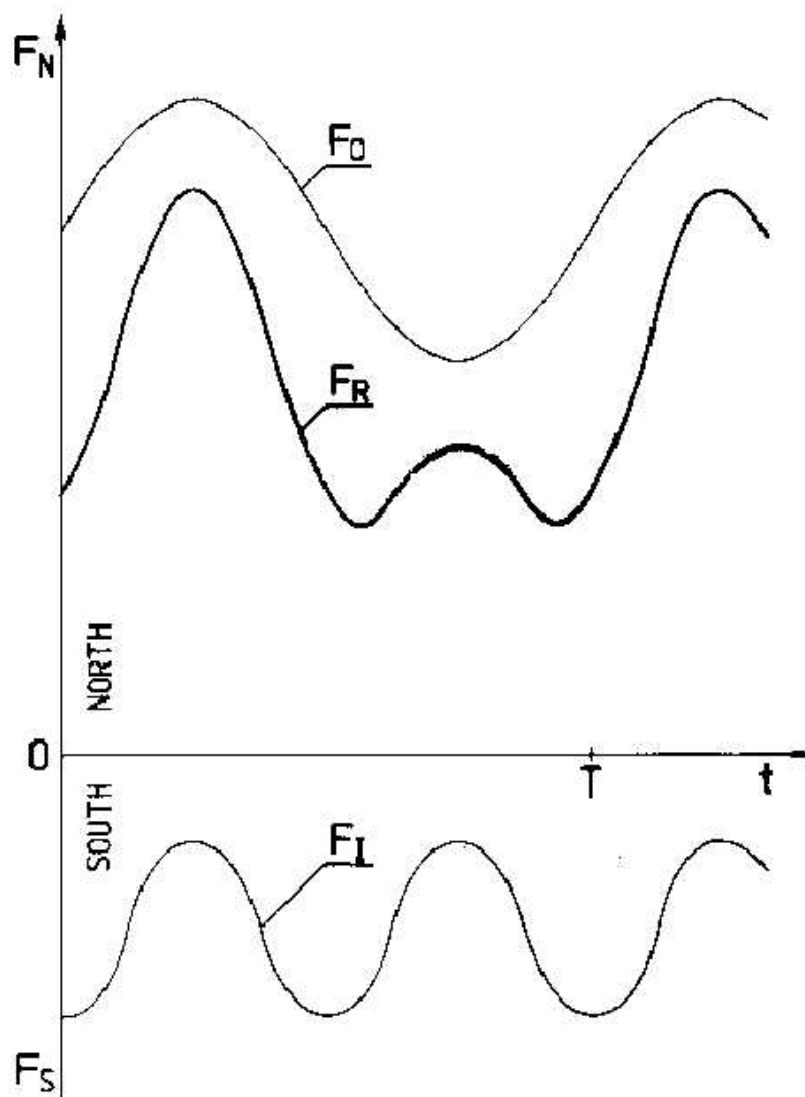


Rys. C6. Wygląd kapsuł dwukomorowych pierwszej generacji w obu trybach pracy.

Zilustrowano tu różnice w wyglądzie zewnętrznym sześciennych kapsuł pierwszej generacji działających w dwóch przeciwnych trybach pracy, tj.: (a) dominacji strumienia WEWNĘTRZNEGO, oraz (b) dominacji strumienia ZEWNĘTRZNEGO. Ponieważ potężne pulsujące pole magnetyczne panujące w takich kapsułach jest przeźroczyste tylko jeśli patrzeć na nie wzdłuż jego linii sił, zakrzywione linie sił strumienia krążącego "C" muszą być nieprzeźroczyste dla postronnego obserwatora, i stąd będą one widoczne jako obszary czerni albo "czarne dziury" (porównaj opis z podrozdziałów F10.3 i F10.4 z rysunkiem C5).

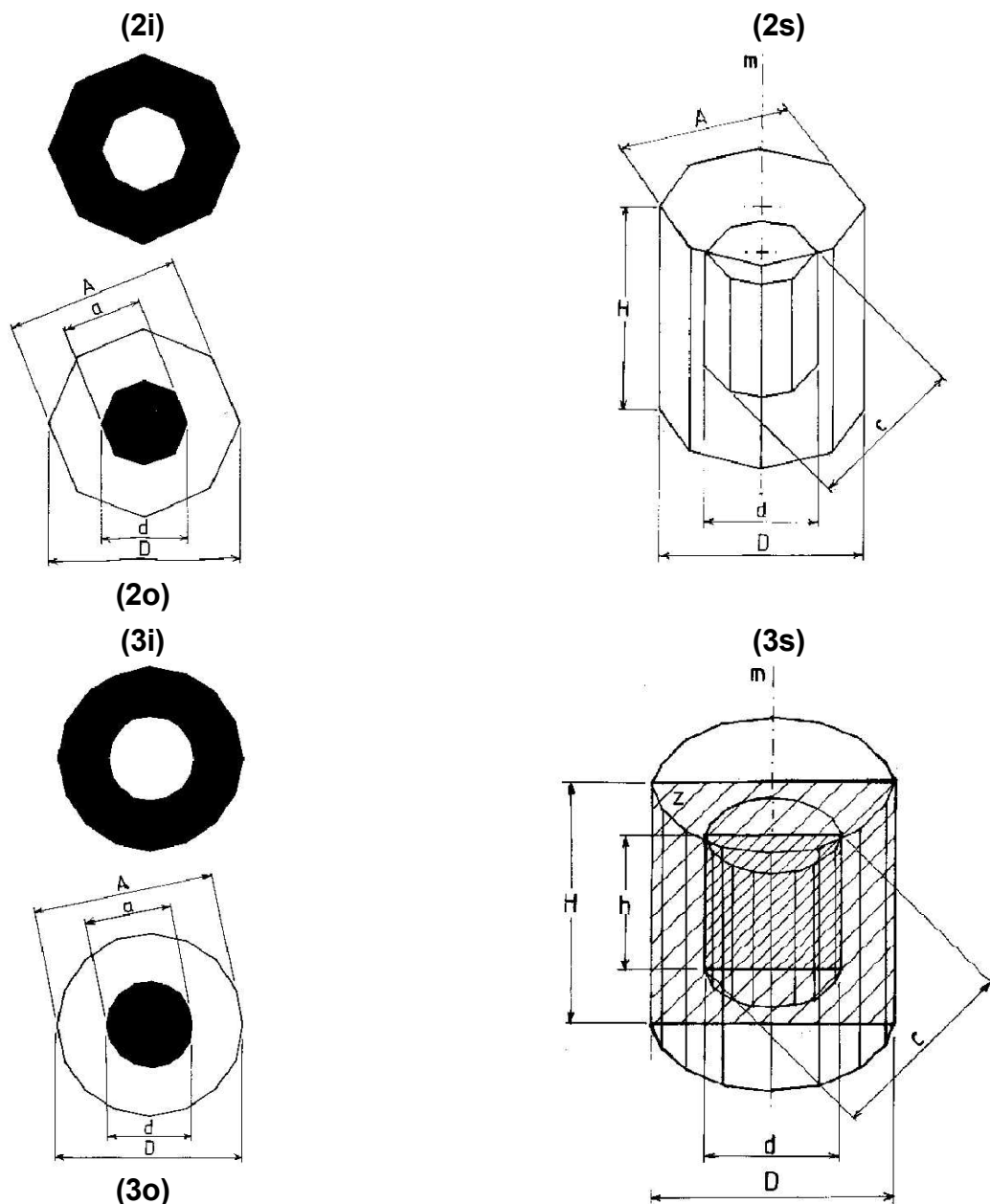
a) Kapsuła pierwszej generacji pracująca z dominacją strumienia wewnętrznego. Strumień wynikowy "R" zostaje w niej wytwarzany przez komorę wewnętrzną "I", podczas gdy cały wydatek komory zewnętrznej "O" zamieniany zostaje w strumień krążący "C". Stąd w kapsule takiej przestrzeń pomiędzy komorą wewnętrzną i zewnętrzną jest nieprzenikalna dla światła i wygląda jak obszar całkowicie zaczerwieniony.

b) Kapsuła pierwszej generacji z dominacją strumienia zewnętrznego. Strumień wynikowy "R" jest w niej wytwarzany przez komorę zewnętrzną "O". Komora wewnętrzna "I" dostarcza jedynie strumienia krążącego "C" jaki w swym obiegu po opuszczeniu komory wewnętrznej w całości zakrzywia się z powrotem do komory zewnętrznej. Stąd w tej kapsule przekrój poprzeczny komory wewnętrznej "I" wygląda jak całkowicie zaczerwieniony.



Rys. C7. Zasada formowania strumienia wynikowego " F_R " w kapsule dwukomorowej z rysunku C5. Zilustrowano przypadki składania razem wydatków z obu komór "O" i "I" takiej kapsuły w celu otrzymania strumienia wynikowego " F_R " jakiego zmiany w czasie odpowiadają tzw. "krzywej dudnienia". Komora zewnętrzna "O" produkuje większy wydatek " F_N " jakiego zmiany w czasie (opisane na jego północnym biegunie "NORTH") reprezentowane są na tym wykresie przy pomocy krzywej " F_O ". Natomiast komora wewnętrzna "I" posiada przeciwnie zorientowane bieguny - patrz rysunek C5 i część (b) na rysunku C6. Stąd w kierunku w którym panuje północny "NORTH" biegun komory zewnętrznej, w komorze wewnętrznej skierowany jest biegun południowy "SOUTH" o wydatku " F_S ". Zmiany w czasie wydatku " F_S " z komory wewnętrznej "I", reprezentowane są przez krzywą " F_I ". Jeśli oba wydatki " F_O " i " F_I " o przeciwnych biegunach zostają zestawione razem, wynikowy strumień " F_R " musi reprezentować algebraiczną różnicę ich wartości $F_R = F_O - F_I$. Różnica ta odprowadzana jest na zewnątrz kapsuły dwukomorowej w postaci właśnie strumienia wynikowego " F_R " ("R" na rys. C5). Cały zaś wydatek " F_I " komory wewnętrznej "I" pozostaje uwięziony we wnętrzu kapsuły w postaci strumienia krążącego "C" jaki cyrkuluje wewnętrznie pomiędzy komorą wewnętrzną i komorą zewnętrzną. (W dalszych rozważaniach kształt wynikowej krzywej dudnienia " F_R " będzie w przybliżeniu reprezentowany przez krzywą zawierającą składową stałą " F_O " i składową pulsującą " ΔF " - patrz rysunki C12 i P18.)

Zauważ, że konfiguracja krzyżowa (rysunek C9) wytwarza strumień wynikowy w sposób niemalże identyczny do zilustrowanego powyżej.



Rys. C8. Kapsuły dwukomorowe drugiej i trzeciej generacji. Ich najważniejszym zastosowaniem są pędniki dyskoidalnego magnokraftu i napędu osobistego drugiej oraz trzeciej generacji. Zilustrowano:

(2s) Widok boczny (side view) kapsuły dwukomorowej drugiej generacji. Jest ona złożona z 2 komór oscylacyjnych o przekroju ośmiobocznym, tj. mniejszej komory wewnętrznej "I" (od angielskiego "Inner" = wewnętrzna) oraz większej od niej komory zewnętrznej "O" (od angielskiego "Outer" = zewnętrzna). Porównaj ten rysunek z rysunkami C5 i C6.

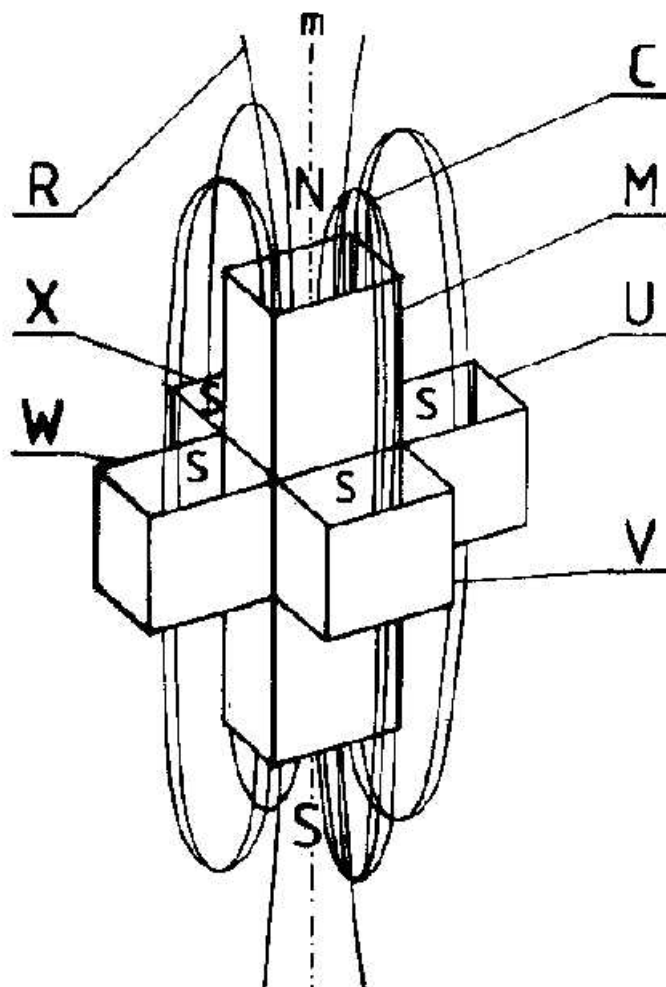
(2i) Widok od góry kapsuły drugiej generacji w trybie dominacji strumienia wewnętrznego.

(2o) Widok od góry kapsuły drugiej generacji pracującej w trybie dominacji strumienia zewnętrznego.

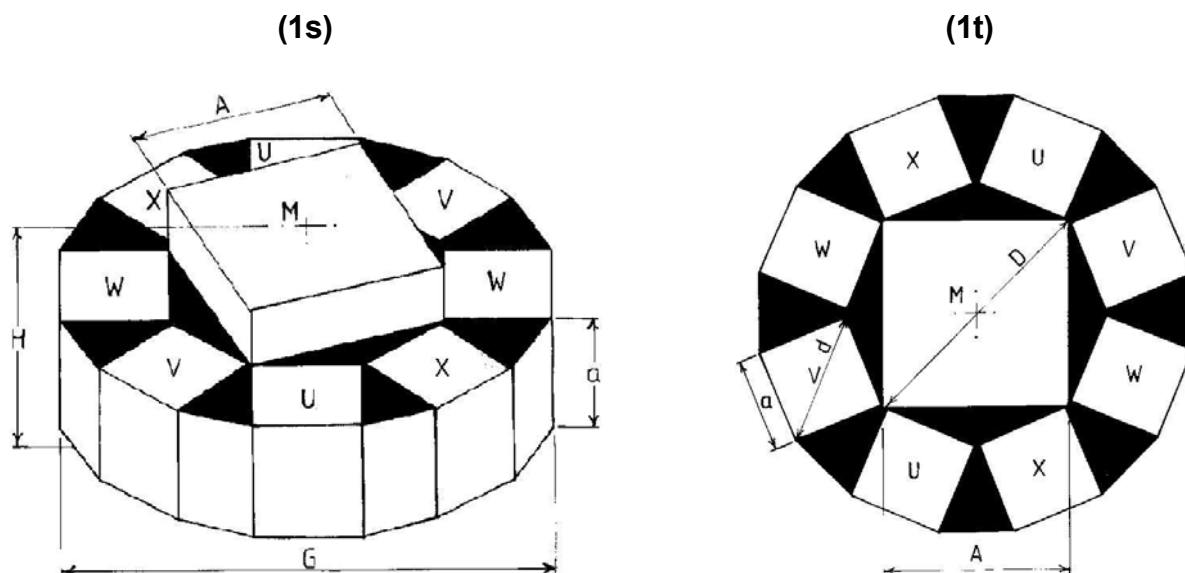
(3s) Widok boczny (side view) kapsuły dwukomorowej trzeciej generacji. Jest ona złożona z 2 komór oscylacyjnych o przekroju 16-bocznym, tj. wewnętrznej (I) oraz zewnętrznej (O).

(3i) Widok od góry kapsuły trzeciej generacji w trybie dominacji strumienia wewnętrznego.

(3o) Widok od góry kapsuły trzeciej generacji w trybie dominacji strumienia zewnętrznego.



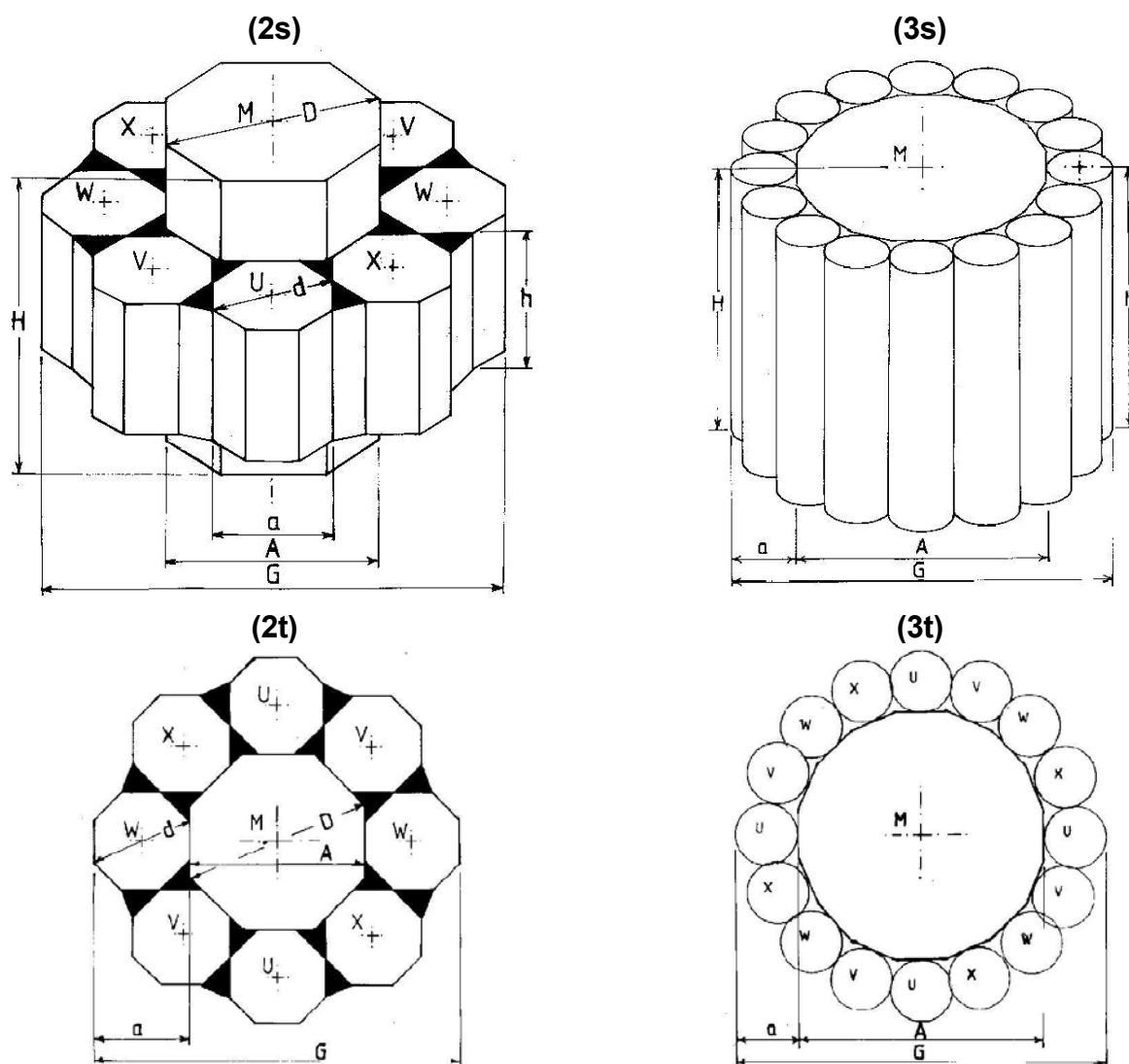
Rys. C9. Standardowa **konfiguracja krzyżowa** pierwszej generacji. Jej najważniejszym zastosowaniem będzie pędnik magnokraftu czteropędnikowego - patrz rysunek D1 (b) i (c). (W początkowym okresie po zbudowaniu danych komór oscylacyjnych może ona także być stosowana w pędnikach dyskoidalnego magnokraftu.) Jest ona uformowana z pięciu komór oscylacyjnych posiadających taki sam przekrój poprzeczny (konfiguracje wyższych generacji mają ich 9 lub 17 - patrz podrozdział C7.2.1). Cztery sześciennie komory boczne (oznaczone jako U, V, W i X) otaczają przeciwstawnie do nich zorientowaną komorę główną (oznaczoną M) jaka jest od nich cztery razy dłuższa. Całkowita objętość wszystkich komór bocznych musi być równa objętości komory głównej. Konfiguracja krzyżowa stanowi uproszczony model układu napędowego magnokraftu. Wynikowy strumień magnetyczny (R) cyrkulowany z niej do otoczenia otrzymuje się jako różnicę pomiędzy wydatkami komory głównej i przeciwstawnie do niej zorientowanych komór bocznych. Zasada formowania tego strumienia wynikowego jest podobna do tej zilustrowanej na rysunku C7. Strumień krążący (C) jest zawsze formowany z wydatku tych komór które wytwarzają mniejszy strumień magnetyczny (w pokazanym na tym rysunku przypadku dominacji strumienia komory głównej - z wydatku wszystkich komór bocznych). Linie sił pola w strumieniu krążącym zawsze zamykają swój obieg poprzez dwie komory. Konfiguracja krzyżowa, podobnie jak kapsuła dwukomorowa, także umożliwia pełne sterowanie wszystkich parametrów wytwarzanego przez nią pola. Jednakże na dodatek do sterowania osiąganego w kapsule dwukomorowej, będzie ona ponadto zdolna do zawirowywania swego pola wokół osi magnetycznej (m), formując w ten sposób własny wir magnetyczny. Jej wadą w porównaniu z kapsułą dwukomorową będzie jednak brak możliwości całkowitego "wygaszenia" pola magnetycznego odprowadzanego przez tą konfigurację do otoczenia (tj. nawet jeśli cały jej wydatek uwięziony zostaje w strumieniu krążącym C, strumień ten ciągle cyrkuluje poprzez otoczenie).



Rys. C10. Prototypowa konfiguracja krzyżowa pierwszej generacji. Jest ona złożona wyłącznie z komór o kształcie sześciątów. Stąd zostanie zbudowana jako nasza pierwsza konfiguracja komór oscylacyjnych poddająca się efektywnemu sterowaniu. Znajdzie się więc w użytkowaniu na długo przed dopracowaniem pierwszej kapsuły dwukomorowej pokazanej na rysunku C5, a także przed dopracowaniem pierwszej standardowej konfiguracji krzyżowej pokazanej na rysunku C9. (Wszakże zbudowanie takiej pierwszej kapsuły dwukomorowej wymagało będzie uprzedniego znalezienia rozwiązania technicznego dla złożonego problemu sterowania swobodnie pływającą komorą wewnętrzną. Z kolei zbudowanie pierwszej standardowej konfiguracji krzyżowej z rys. C9 wymagało będzie dopracowania komory głównej (M) o długości czterokrotnie przewyższającej szerokość jej boków.) W początkowej fazie budowy naszych wehikułów z napędem magnetycznym pokazana powyżej prototypowa konfiguracja krzyżowa będzie montowana nawet w pędnikach dyskoidalnego magnokraftu - patrz etap (1A) w klasyfikacji z podrozdziału M6. Zasada działania tej konfiguracji prototypowej jest identyczna do zasady działania standardowej konfiguracji z rysunku C9. Jedyna różnica sprowadza się do formowania dwóch fal magnetycznych zamiast jednej. Konfiguracja prototypowa jest łatwa do rozpoznania po swoim dyskoidalnym kształcie w którym jej szerokość $G = 2A$ jest dwa razy większa od wysokości $H = A$. Zilustrowane wymiary obejmują: $A = 2a$ - długość boku komory głównej (M), $a = (\frac{1}{2})A$ - długość boku komór bocznych (U, V, W i X), $H = A$ - wysokość całej konfiguracji, D , d - średnice okręgów opisanych na czołach komory głównej i komór bocznych.

(1s) Widok boczny (side view) całej konfiguracji. Zaczerniono tworzywo nośne.

(1t) Widok odgórny (top view) tej konfiguracji.



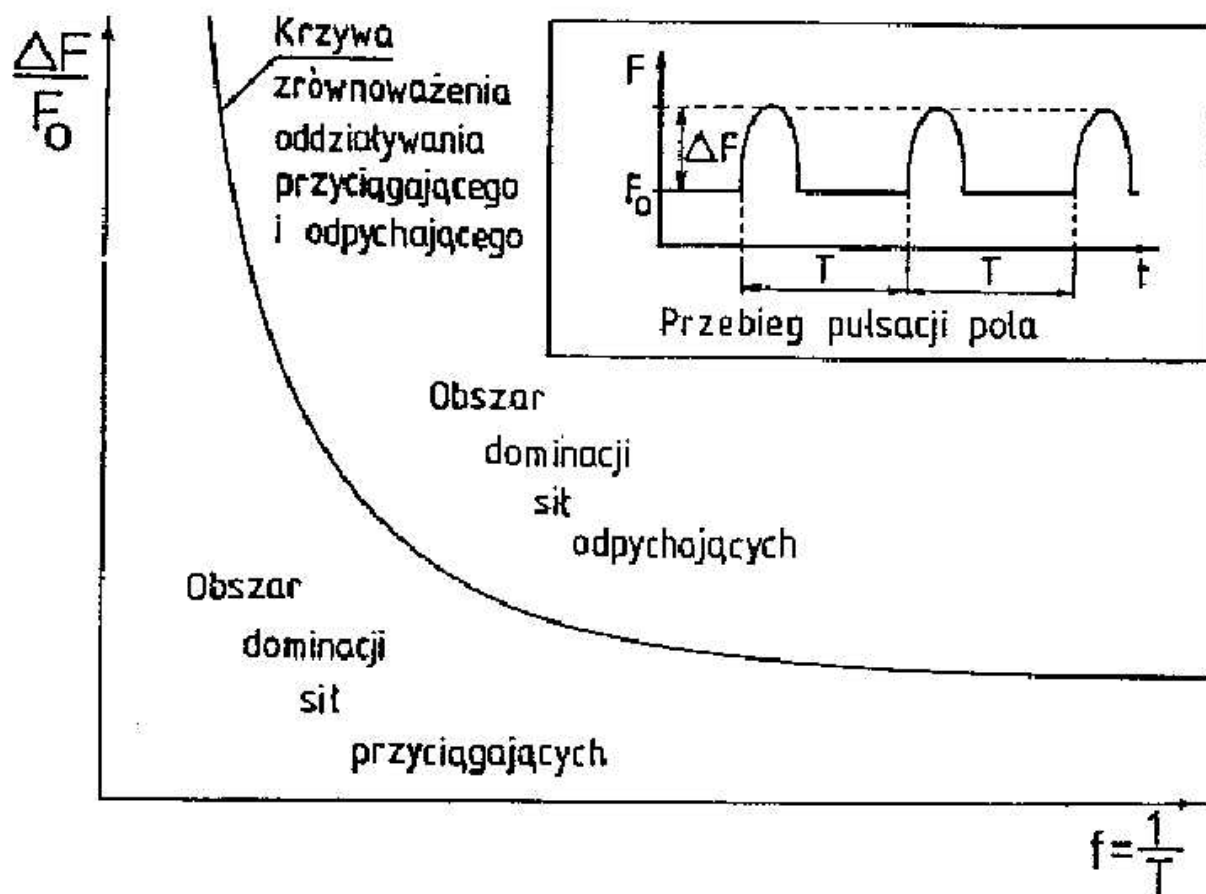
Rys. C11. Konfiguracje krzyżowe drugiej i trzeciej generacji. Ich najważniejszym zastosowaniem będzie pędnik magnokraftu czteropędnikowego drugiej i trzeciej generacji (rysunek D1). (W początkowym okresie po swym zbudowaniu mogą one także być stosowane w pędnikach dyskoidealnego magnokraftu.) Zinterpretowano wymiary: A, a, D, d, H, h, G. Rysunek pokazuje:

(2t) Widok z góry (top view) konfiguracji krzyżowej drugiej generacji. Jest ona złożona z 9 komór oscylacyjnych o przekroju ośmiobocznym, tj. jednej komory głównej (M) oraz ośmiu takich komór bocznych (U, V, W i X) formujących dwie fale magnetyczne. Zaczerniono tworzywo nośne.

(2s) Widok boczny (side view) konfiguracji krzyżowej drugiej generacji (porównaj z rys. C10).

(3t) Widok z góry (top view) konfiguracji krzyżowej trzeciej generacji. Jest ona uformowana z 17 komór oscylacyjnych o przekroju szesnastobocznym, tj. jednej komory głównej (M) oraz szesnastu komór bocznych (U, V, W i X) formujących cztery fale magnetyczne. Zauważ że dla niej $A = 4a$.

(3s) Widok boczny (side view) konfiguracji krzyżowej trzeciej generacji. Zauważ że dla niej $H=h$.



Rys. C12. "Krzywa równowagi" oddziaływań pomiędzy polem magnetycznym wytwarzanym przez kapsułę dwukomorową lub konfigurację krzyżową, a przedmiotami ferromagnetycznymi zawartymi w ich otoczeniu. Jak to powszechnie wiadomo, stałe pole magnetyczne przyciąga przedmioty ferromagnetyczne. Stąd też wszystkie pola w jakich **składowa stała (F_0) dominuje** nad składową pulsującą (ΔF) muszą przyciągać przedmioty ferromagnetyczne. Parametry pola w którym ta składowa stała przewyższa składową zmienną znajdują się poniżej krzywej z tego wykresu (tj. w obszarze dominacji sił przyciągających). Eksperymenty z polami dynamicznymi ustaliły, że pulsujące pole magnetyczne wypycha ze swego zasięgu wszystkie przedmioty przewodzące (w więc także ferromagnetyki). Stąd też wydatki kapsuły dla których **składowa pulsująca (ΔF) dominuje** nad składową stałą (F_0) będą powodować odpychanie wszystkich przedmiotów ferromagnetycznych. Pola magnetyczne w których składowa pulsująca (ΔF) dominuje nad składową stałą (F_0) leżą ponad krzywą z tego wykresu (tj. w obszarze dominacji sił odpychających). Natomiast dla parametrów pola magnetycznego w którym **obie składowe balansują** swoje działanie, tj. leżących dokładnie na pokazanej tu krzywej równowagi, przyciąganie i odpychanie nawzajem się zrównoważą. Stąd pole jakie oznacza się takimi parametrami nie będzie ani przyciągało ani też odpychało przedmiotów ferromagnetycznych zawartych w jego zasięgu. Pole takie będzie więc się zachowywało jak rodzaj jakiejś hipotetycznej "antygravitacji" raczej niż pola magnetycznego.

Obramowanie zawiera interpretację wszystkich dyskutowanych parametrów pulsującego pola objaśnianego na tym rysunku.



Rys. C13. Przykład stanowiska do badań nad komorą oscylacyjną. Zostało ono opracowane, zbudowane i przebadane przez hobbystę z Bydgoszczy który jednak prosił aby nie publikować jego nazwiska i adresu.

(Góra) Model komory oscylacyjnej sfotografowany w ciemności. Ukazuje on fascynujący wygląd pęków rotujących iskier elektrycznych. Fotografia ta wykonana została w maju 1987 roku.

(Dół) Fotografia eksperymentalnego stanowiska badawczego składającego się z: badanego prototypu komory, generatora impulsów jaki dostarcza komorze jej mocy elektrycznej, elektromagnesu odchylającego, oraz urządzeń pomiarowych. Sfotografowano w sierpniu 1989 roku.

Zastosowania komory oscylacyjnej

Jak dotychczas nie istnieje żaden inny wynalazek jaki zmieniłby stan naszego otoczenia technicznego do tego samego stopnia jak to uczyni skompletowanie komory oscylacyjnej. Impakt jaki to urządzenie będzie posiadało na aspekt materialny naszego życia może być jedynie porównany do efektu w intelektualnej sferze wywołanego wprowadzeniem tam komputerów. Istnieje wysokie prawdopodobieństwo, iż do około 2084 roku (tj. w sto lat po wynalezieniu komory oscylacyjnej) prawie każde aktywne urządzenie wykorzystywane przez ludzi będzie zawierało jakąś formę komory oscylacyjnej. Wiele obiektów które obecnie są pasywne, takich jak meble, budynki, monumenty, itp., będzie przetransformowane przez komorę oscylacyjną w struktury aktywne, tj. będą one się poruszały, zmieniały swoje zorientowanie i dopasowywały swoje położenie do zmieniających się wymagań ich użytkowników. Przeglądnijmy więc pokrótce najważniejsze zastosowania komory oscylacyjnej, starając się przewidzieć jaki wpływ to urządzenie będzie posiadało na dany obszar ludzkiej działalności.

Najsilniejszy impakt posiadało będzie wprowadzenie komory oscylacyjnej do **energetyki**. Praktycznie przetransformuje ona kompletnie obecne metody wytwarzania, przesyłania i konsumowania energii. Ogromna różnorodność odmiennych urządzeń jakie obecnie wykorzystywane są w tym celu, po pojawieniu się komory zastąpiona zostanie przez jeden rodzaj uniwersalnej kapsuły dwukomorowej po odpowiednim przesterowaniu zdolnej do wypełniania setek funkcji. Aby uzmysłowić jak ogromnemu przeobrażeniu ulegnie wówczas obraz naszej planety, wystarczy tu wspomnieć iż przykładowo wszystkie obecne linie przesyłowe (wysokiego i niskiego napięcia) całkowicie znikną ponieważ energia rozprzestrzeniana będzie po upakowaniu jej w "konserwy", tj. niewielkie, lekkie, poręczne, oraz ponownie ładowalne kapsuły dwukomorowe. Z kolei szerokie użycie komór oscylacyjnych zamiast obecnych linii przesyłowych ogromnie poprawi stronę zdrowotną pól otaczających naszą planetę. Wydzielane bowiem przez dzisiejsze linie przesyłowe pola elektryczne i elektromagnetyczne będą niemal całkowicie wyeliminowane. Ponadto częstość robocza każdej kapsuły dwukomorowej tak może zostać dobrana aby produkowała ona wyłącznie stymulujące zdrowie i sprawdzające dobre samopoczucie wibracje telepatyczne - patrz podrozdziały H7.1, NB2, NB3.

Nowe horyzonty w **wytwarzaniu** i dostarczaniu energii otworzy wykorzystanie wielowymiarowej transformacji energii zachodzącej w komorze oscylacyjnej. W jej efekcie przewidzieć można zastąpienie układami komór wszelkich obecnych urządzeń jakie służą produkcji lub transformacji energii. I tak obecne silniki spalinowe, generatory, ogniwa foto- lub termo-elektryczne, transformatory, itp., wszystkie one przyjmą formę kapsuł dwukomorowych - patrz tablica C1. Z uwagi na ich wysoką sprawność (tj. pracę praktycznie bez strat energii), dostarczą one naszej cywilizacji wymaganej przez nią energii bardziej efektywnie oraz w sposób mniej szkodliwy dla naturalnego środowiska.

Komora oscylacyjna umożliwi też opracowanie i szerokie wprowadzenie nowych, środowiskowo bardziej **"czystych" metod** wytwarzania energii. Takie urządzenia jak telekinetyczne urządzenia do pozyskiwania energii otoczenia (opisane w rozdziale K i monografii [6]) oraz generatory czystej energii (wykorzystujące promieniowanie słoneczne, wiatr, fale oceaniczne, przyływy i odpływy morza, itp.), staną się niezwykle efektywne jeśli oparte zostaną na wykorzystaniu kapsuł dwukomorowych.

Liczne energetyczne zastosowania komory oscylacyjnej wynikną w przyszłości z jej zdolności do **akumulowania** ogromnych ilości energii. Aby dać nam przedsmak potencjału jaki to urządzenie kryje w sobie, wystarczy wspomnieć iż zapotrzebowania energetyczne

współczesnej fabryki, miasta, dużego okrętu czy samolotu, mogą zostać zaspokajane komorą o wielkości główki od szpilki - jeśli tylko będziemy w stanie zbudować ją aż w tak małych wymiarach. Wszystkie więc obecne baterie, akumulatory, oraz generatory awaryjne, zastąpione zostaną przez efektywne i ponownie ładowalne komory oscylacyjne. Budowane jako kapsuły dwukomorowe, w przypadku takiego użycia jako akumulatory energii, nie będą one odprowadzały do otoczenia żadnego pola magnetycznego.

Prawie wszystkie obecne urządzenia **transformujące** energię, przykładowo latarnie, systemy oświetleniowe ulic i pomieszczeń, grzejniki, klimatyzatory powietrza, silniki elektryczne, itp., zastąpione zostaną przez odpowiednie nasterowanie odmiennych funkcji u tych samych kapsuł dwukomorowych.

Dzięki komorze oscylacyjnej transformacja energii w przyszłości zastąpi również obecną **transformację ruchu**. Stąd przyszłe mechanizmy będą znacznie prostsze i lżejsze, ponieważ zostaną one uwolnione od zawierania w sobie wszystkich tych dodatkowych urządzeń jakie obecnie dostarczają i transformują ruch. W przyszłości ruch będzie wytwarzany w dokładnym miejscu gdzie zachodzi jego spożytkowanie, a także i w dokładnej formie w jakiej jest on wymagany. Dla przykładu, jeśli w przyszłości jakiś hobbysta zechce zbudować kopię naszego dzisiejszego samochodu, wyprodukuje on ruch we wnętrzu kół poprzez wstawienie tam kilku kapsuł dwukomorowych. Stąd cały dzisiejszy silnik, skrzynia biegów, oraz transmisja staną się niepotrzebne.

Unikalne zalety komory oscylacyjnej spowodują, że to urządzenie całkowicie przejmie obecne funkcje **elektromagnesów**. Laboratoria badawcze, zdolne do użycia precyzyjnie sterowalnych pól magnetycznych o obecnie nieosiągalnej mocy (a także przebiegu ich zmian czasowych - np. pól pulsujących desymetrycznie czy pól stałych), będą zdolne do wydarcia naturze wielu sekretów, wprowadzając w ten sposób ogromny postęp do naszej nauki i techniki. Przemysł, wykorzystując technologie jakie będą bazowały na wykorzystaniu supersilnych pól magnetycznych, dostarczy ludziom wielu produktów dotychczas jeszcze niemożliwych do wytworzenia. Dla przykładu, przemysł ten wyprodukować może niezniszczalną gumę i odzież, obiekty w całości wykonane z monokryształów, beton silniejszy od stali, itp. Także nowy rodzaj mognetoreflexyjnego materiału, zdolnego do wypełnienia wymagań magnetycznych komory oscylacyjnej, wyprze te znajdujące się w użyciu obecnie.

Komora oscylacyjna nie tylko wyeliminuje elektromagnesy stosowane jako oddzielne urządzenia, ale także te jakie wchodzą w skład innych urządzeń jako ich podzespoły, np. z silników elektrycznych, generatorów elektryczności, itp. Zalety komory, takie jak: wysoki stosunek mocy-do-wymiarów, zdolność do znoszenia długich przerw pomiędzy chwilą dostarczenia energii i czasem użycia tej energii, sterowalność; wynikną w szerokim użyciu tego urządzenia do budowy lekkich wehikułów, pomp i generatorów pracujących daleko od źródeł energii i centrów cywilizacyjnych, silników okrętowych i lotniczych, itp.

Kapsuły dwukomorowe dostarczające **stałego pola** magnetycznego zastąpią też dzisiejsze magnesy stałe. Stąd przyszłe modele naszych głośników, łożysk, sprzęgieł, chwytaaków, szyn, itp., wszystkie one wykorzystywały będą komory oscylacyjne.

W przyszłości komory oscylacyjne modulowane sygnałami myślowymi będą też stanowić niezwykle efektywne urządzenia do **łączości** telepatycznej, umożliwiające swym użytkownikom natychmiastowe komunikowanie się z najodleglejszymi zakątkami naszego wszechświata (patrz opisy w rozdziale N).

Komora oscylacyjna wprowadzi także zupełnie nową **modę**, jaka w dzisiejszych czasach nie posiada odpowiedniego zabezpieczenia technicznego. Będzie to moda na zawieszanie obiektów w przestrzeni. Należy więc się spodziewać, iż przyszłe meble, urządzenia domowe, maszyny wytwórcze, a nawet całe budynki i elementy architektoniczne, będą wisiały w przestrzeni, podtrzymywane przez niewidzialne linie sił pola magnetycznego. Dla przykładu taki mebel jak dzisiejszy fotel, w przyszłości będzie szybował po przestrzeni mieszkania, zaś wbudowany w niego komputer będzie analizował ustne lub myślowe polecenia siedzącej na nim osoby, przenosząc tą osobę we wymagane miejsce, zmieniając jej orientację, wysokość i nachylenie, a także adaptując swój kształt do typu postawy

wypoczynkowej jaką ta osoba zapragnie w danej chwili przyjąć. Jedną z konsekwencji tej mody na zawieszanie obiektów w przestrzeni będzie niemal całkowite zaniknięcie koła, jako iż obecne ruchy toczące zostaną zastąpione przez szybowanie.

Oczywiście ogromny potencjał kryje się w **militarnym** użyciu komory oscylacyjnej. Może ona zarówno zwielokrotnić możliwości już istniejących urządzeń i środków bojowych, jak i uformować dotychczas jeszcze nie znane rodzaje broni. Aby zilustrować potencjał komory w zwielokrotnianiu możliwości już istniejących rodzajów broni wystarczy wspomnieć iż ilość energii zakumulowana w kapsule dwukomorowej wielkości kostki do gry wystarcza będzie aby utrzymać bombowiec w powietrzu przez całe lata bez konieczności jego lądowania w celu ponownego zatankowania, aby przepłynąć łodzią podwodną w stanie zanurzenia kilkaset razy naokoło naszego globu, czy aby przejechać czołgiem drogę większą od odległości Ziemi od Słońca. Aby ukazać potencjał komory oscylacyjnej w formowaniu nowych rodzajów broni, wystarczy tu wspomnieć iż układ tych urządzeń wytwarzający wirujące pole magnetyczne będzie w stanie uformować zapory i pola minowe jakie w ciągu sekund mogą odparować eksplozyjnie każdy obiekt wykonany z dobrego przewodnika elektryczności jaki wejdzie w ich obszar działania. Pociski zawierające układy komór z takim wirującym polem, mogą spowodować natychmiastowe wyparowanie ogromnych konstrukcji wykonanych ze stali, takich jak mosty, fabryki, okręty, samoloty, rakiety, satelity, itp. Z kolei gwałtowne uwolnienie ogromnej energii zgromadzonej w komorze (np. poprzez jej zdetonowanie - patrz podrozdział F14 lub monografia [5]) spowoduje eksplozję porównywalną w efektach do użycia bomby termojądrowej. Jedyną różnicą będzie, iż po eksplozji komory otoczenie nie zostanie skażone radioaktywnością, stąd będzie się nadawało do natychmiastowego zajęcia i ponownego zasiedlenia. Z uwagi przy tym na niewielkie rozmiary komór, potencjał do formowania zniszczeń odpowiadających wybuchowi sporej bomby termojądrowej uzyska mała kapsuła dwukomorowa mieszcząca się w zwykłym pocisku karabinowym. Oczywiście komory oscylacyjne nie tylko są w stanie niszczyć, ale umożliwiają też osłanianie się przed zostaniem zniszczonym przez przeciwnika. Najprostsza taka osłona polegała będzie na zaopatrzeniu wybranych wehikułów lub obiektów wojskowych w komory oscylacyjne których pola będą formowały odpychające lub przyciągające oddziaływania z obiektami ferromagnetycznymi ze swego otoczenia (patrz rysunek C12). W ten sposób będą one w stanie odepchnąć (lub - w razie konieczności, także przyciągnąć, obezwładnić i przechwycić) dowolne wehikuly i pociski strony przeciwnej. Bardziej niezwykła możliwość komór oscylacyjnych wynika z możliwości formowania przez nie tzw. "soczewki magnetycznej" (opis tej soczewki zawarty został w podrozdziale F10.3). Osłonięci nią żołnierze, wehikuly, lub obiekty o znaczeniu militarnym staną się całkowicie niewidzialni/e dla przeciwnika.

Najbardziej jednak zachęcające perspektywy otwiera użycie komory oscylacyjnej do przeznaczenia dla którego jej zasada została oryginalnie wynaleziona, tj. do celów transportowych. Przy takim jej użyciu, najważniejsze jej zastosowanie polegać będzie na pełnieniu funkcji urządzenia **napędowego** (tj. pędnika) dla napędów osobistych, wehikułów latających, oraz statków międzygwiezdnych. Z upływem czasu wypracowane także będzie transportowe użycie komór oscylacyjnych w tzw. "urządzeniach zdalnego oddziaływania", których przykładami może być odpowiednio nasterowane "pole podnoszące" kapsuł dwukomorowych opisane w podrozdziale C7.3, czy tzw. telekinetyczny "promień podnoszący" opisany w podrozdziale H6.2.1. Rozdziały F, D, E, L i M niniejszej monografii poświęcone zostały szerszemu omówieniu transportowych zastosowań komory jako pędnika dla wehikułów latających.

Na zakończenie przytoczonego tu przeglądu zastosowań komory warto podkreślić iż wszystkie funkcje opisane w tym podrozdziale wypełniane mogą być przez tą samą kapsułę dwukomorową zaopatrzoną jedynie w odmienny system/program sterowania. Stąd w sensie uniwersalności swych zastosowań komory oscylacyjne przypominać będą współczesne komputery w których jedynie zmiana programu sterującego przekształca je przykładowo z maszyny do pisania w instrument muzyczny, automatycznego pilota, atlas drogowy, kasyno gier, czy przyrząd pomiarowy.

Do tego miejsca przegląd zastosowań komór oscylacyjnych dokonywany został w odniesieniu do pierwszej generacji tych urządzeń. Niemniej - jak to czytelnik zapewne odnotował już z podrozdziału C4.1, po zrealizowaniu komór pierwszej generacji, na Ziemi podjęty zostanie rozwój komór drugiej, a później nawet trzeciej generacji. Komory tych wyższych generacji użyteczne będą w tych wszystkich opisywanych tutaj zastosowaniach co komory pierwszej generacji. Jednak na dodatek to powyższego będą one wypełniały dodatkowe funkcje którym komory pierwszej generacji nie są w stanie sprostać. Przykładowo komory drugiej generacji wytwarzały będą pole telekinetyczne o ogromnej aktywności biologicznej. Stąd będzie je można dodatkowo wykorzystywać jako maszyny leczące (patrz opisy w podrozdziałach N5.2, HB3 i T1), jako urządzenia wzbudzające płodność u kobiet (podrozdział NB3), czy też jako źródło stałego pola telekinetycznego użytecznego do telekinetyzowania środowiska w telekinetycznym rolnictwie - patrz opisy w podrozdziale NB2 niniejszej monografii (a także w podrozdziale G2.1.1.2 monografii [5]). W podobny sposób ich pole telekinetyczne może też służyć jako katalizator trudnych do przeprowadzenia reakcji chemicznych lub modyfikator własności materiałów (podrozdział NB3), czy jako nośnik informacji telepatycznej (podrozdziały H7.1 i N1).

Niezależnie od zastosowaniowego znaczenia komory oscylacyjnej, zbudowanie tego urządzenia będzie także posiadało ogromne **znaczenie poznawcze**. Komora oscylacyjna będzie bowiem pierwszym "rezonatorem magnetycznym" zbudowanym na naszej planecie jaki efektywnie wytwarzał będzie własne drgania magnetyczne a także posiadał będzie zdolność do reagowania na takie drgania pochodzące z innych źródeł. Aczkolwiek nauka ziemską stoi dopiero na początku swej drogi do poznania możliwości i znaczenia drgań magnetycznych, już obecnie wiadomo iż stanowią one klucz do ogromnej ilości dotychczas nieopanowanych jeszcze zjawisk, do których przykładowo zaliczyć można opisane w rozdziale M podróże w czasie i telekinezę czy postulowane Konceptem Dipolarnej Grawitacji: telepatię, zdalne kontrolowanie psychiki ludzkiej i nastrojów społecznych (po więcej szczegółów patrz podrozdziały N4, T5 i V4.2, w tej monografii oraz D4 w monografii [5]), uzdrawianie, transformowanie jednych pierwiastków i substancji w inne, pozyskiwanie energii otoczenia, oraz wiele innych. Stąd w sensie poznawczym komora oscylacyjna stanowić będzie prototyp i poprzednika dla całej gamy nadchodzących po niej urządzeń wytwarzających, przetwarzających, wykrywających i mierzących drgania magnetyczne, przyczyniając się w ten sposób do uformowania w przyszłości całych nowych dziedzin nauki i techniki. Dla dalszych generacji naukowców i inżynierów na Ziemi jej znaczenie poznawcze prawdopodobnie będzie równie przełomowe jak znaczenie obwodu Henry'ego było dla dzisiejszych elektroników czy cybernetyków.

CB1. Przyszłe zastosowania komory oscylacyjnej jako akumulatora do bezspalinowych samochodów (czyli do tzw. „eco-cars”)

The Oscillatory Chamber is able to accumulate in relatively small space theoretically unlimited amounts of energy. So in practice in a size similar to present car batteries, such an Oscillatory Chamber is able to store the amount of energy that would suffice for several thousands of years of use of present cars. As such these "Oscillatory Chambers" can effectively replace batteries from present electric cars.

CB2. Senator McCain obiecał nagrodę w wysokości 300 millions dolarów dla wynalazcy akumulatora energii który by wykazywał cechy komory oscylacyjnej

Kandydat na prezydenta USA, senator John McCain, we wtorek dnia 24 czerwca 2008 roku obiecał publicznie że ufunduje nagrodę w wysokości 300 milionów dolarów USA temu

wynalazcy który wynajdzie korzystny dla naturalnego środowiska akumulator energii nowej generacji, przydatny do napędzania samochodów. Jego obietnica została natychmiast rozgłoszona po świecie. Już następnego dnia powtarzały ją niemal wszystkie dzienniki telewizyjne oraz wiele gazet. Przykładowo, w Nowej Zelandii opisana ona została w artykule "McCain offers \$394m for greener car battery" (tj. "Senator McCain obiecuje \$394 milionów dolarów nowozelandzkich za 'zielony' akumulator do samochodu"), ze strony B1 gazety „Dominion Post”, wydanie ze środy (Wednesday), June 25, 2008. W następnym tygodniu tamta obietnica była komentowana w artykule "Bravo to those extending the knowledge frontiers" (tj. "Brawo dla tych co poszerzają czołówkę wiedzy") ze strony B5 nowozelandzkiej gazety „Dominion Post”, wydanie z wtorku (Tuesday), July 1, 2008.

"Komora oscylacyjna" opisywana w niniejszym rozdziale i monografii jest właśnie owym cudownym akumulatorem nowej generacji w przyszłości przydatnym m.in. i do napędzania samochodów elektrycznych. Wszakże ilość energii którą jest ona w stanie zakumulować w objętości dzisiejszego akumulatora samochodowego, potem zaś uwolnić na życzenie, wystarcza do napędzania samochodu przez kilka tysięcy lat. Ponadto nie generuje ona żadnych zanieczyszczeń. Z kolei energia jest w niej przechowywana w najczystszej formie pulsującego pola magnetycznego które umożliwia łatwe czerpanie tej energii za pośrednictwem zwykłego uzwojenia transformatorowego i następne jej bezpośrednie użycie do napędzania samochodów.

Oczywiście, opisywane tutaj obwieszczenie senatora McCain ma wartość głównie jako moralne (a nie finansowe) wsparcie dla badań i rozwoju nad komorą oscylacyjną. Wszakże narazie jest ono tylko obietnicą, a nie faktycznym dofinansowaniem. Niemniej nawet będąc jedynie obietnicą, ciągle ma ono dużą wartość jako uwypuklenie wagi i pilności technicznego urzeczywistniania idei komory oscylacyjnej. Wszakże uzmysławia ono każdemu że rozwój sytuacji z zasobami ropy naftowej na Ziemi nieodwołalnie wiedzie do sytuacji że któregoś dnia "komora oscylacyjna" stanie się ludzkości absolutnie niezbędna. Dzień zaś taki jest już coraz bliżej. W owym zaś krytycznym czasie na wagę złota będzie ekspertyza badaczy którzy mają już jakieś doświadczenie w badaniach i w rozwoju "komory oscylacyjnej". Dlatego ja osobiście bym rekomendował, aby każdy kto ma dostęp do odpowiedniej prototypowni oraz do możliwości prowadzenia badań laboratoryjnych, już obecnie włączył się do budowy i rozwoju "komory oscylacyjnej". Zainwestowanie jego zainteresowań w ów niezwykle akumulator energii bezsprzecznie pewnego dnia musi okazać się ogromnie korzystne.

MAGNOKRAFT CZTEROPĘDNIKOWY

Motto: "Kiedy zobaczysz w powietrzu coś co wygląda jak stodoła, nie ograniczaj wyjaśnień do fermentującego siana, romansu myszy z nietoperzami, żartu studentów, czy pijanej muchy na okularach, a rozważ również iż mógł to być nie znany ci jeszcze statek kosmiczny."

Niniejszy rozdział zaprezentuje konstrukcyjnie i działaniowo najprostrzy wehikuł którego napęd wykorzystuje komory oscylacyjne. (Niestety, pod względem wykonawczym i sterowania jest on wehikułem najtrudniejszym do zbudowania.) Statkowi temu przyporządkowana została nazwa **"magnokraft czteropędnikowy"** albo "wehikuł czteropędnikowy". Aby wyraźnie odróżnić ów magnokraft czteropędnikowy od statku opisanego w rozdziale F, tamten omówiony w rozdziale F statek w niniejszej monografii okreśłany będzie jako **"dyskoidalny magnokraft"** lub po prostu **"magnokraft"**. Wehikuł czteropędnikowy, obok dyskoidalnego magnokraftu, reprezentuje jedno z dwóch podstawowych zastosowań pędników magnetycznych. (Trzecie zastosowanie tych pędników, które stanowi jedynie nieco zmodyfikowaną wersję dyskoidalnego magnokraftu, nazywane jest "magnetycznym napędem osobistym" i opisane zostało w rozdziale E.) Podczas jednak gdy działanie systemu napędowego magnokraftu jest najoptymalniejsze jeśli wykorzystuje ono tzw. "kapsułę dwukomorową", zrealizowanie wehikułu czteropędnikowego bezwzględnie wymaga użycia odmiennej niż kapsuła dwukomorowa konfiguracji komór oscylacyjnych, zwanej tu "konfiguracją krzyżową" - patrz opis w podrozdziale C7.2 niniejszej monografii. Każdy pędnik magnokraftu czteropędnikowego zawiera w sobie jedną konfigurację krzyżową. Pole magnetyczne wytwarzane przez tą konfigurację wykazuje obecność wszystkich atrybutów wymaganych dla lotów i manewrowania statku kosmicznego. To z kolei jest przyczyną dla której magnokraft czteropędnikowy może ograniczyć swój system napędowy do jedynie czterech pędników (w przeciwieństwie do minimum ośmiu pędników bocznych plus jeden pędnik główny wymaganych dla lotów dyskoidalnego magnokraftu). Ponieważ liczba pędników jest najbardziej charakterystycznym elementem wehikułu czteropędnikowego, jego nazwa wyraża sobą tą liczbę. Każdy z czterech pędników tego wehikułu zamocowany jest do jednej z czterech naroży kabiny załogi. Stąd, cztery beczko-podobne pędniki wystające na zewnątrz głównego korpusu tego wehikułu dostarczają dodatkowego szczegółu identyfikacyjnego, bardzo charakterystycznego dla tych statków kosmicznych.

Jeśli chodzi o chronologię budowy poszczególnych magnokraftów, wehikuł czteropędnikowy najprawdopodobniej będzie zbudowany jako trzeci wehikuł wykorzystujący komory oscylacyjne (po dyskoidalnym magnokrafcie bazującym na konfiguracjach krzyżowych, oraz dyskoidalnym magnokrafcie bazującym na kapsułach dwukomorowych - patrz okres 1C klasyfikacji podanej w podrozdziale M6). Przyczyną dla tego stanu rzeczy będzie, iż wehikuł ten wymaga znacznie bardziej wyrafinowanych systemów sterowania niż dyskoidalny magnokraft opisywany w rozdziale F. Systemy te będą więc mogły zostać opracowane dopiero kiedy w efekcie eksploatacji dyskoidalnego magnokraftu nasza cywilizacja zakumuluje odpowiedni zasób wiedzy o sterowaniu statków z napędem magnetycznym. Aczkolwiek więc konfiguracje krzyżowe stosowane w pędnikach wehikułu czteropędnikowego są prostsze w budowie niż kapsuły dwukomorowe stosowane w pędnikach dyskoidalnego magnokraftu, owe zaostrome wymagania na system sterowniczy spowodują, iż wehikuł ten będzie musiał odczekać nieco na swoją kolejność do zbudowania (patrz podrozdział M6).

Ogólna budowa (konstrukcja) i wygląd magnokraftu czteropędnikowego pokazane zostały w części (a) **rysunku D1**. Wehikuł ten składa się z dwóch zasadniczych podzespołów, tj.: korpusu (2) i pędników (3).

Korpus główny (2) stanowi zasadniczy element magnokraftu czteropędnikowego. Korpus ten najczęściej przyjmie formę sześciennego lub prostopadłościennego domku czy stodoły. Na wierzch tego domku/stodoły nałożony zostaje dach (1) w kształcie piramidki, jaki nadaje wynikowej konstrukcji wymaganej aerodynamiczności, a jednocześnie za pośrednictwem swoich proporcji wymiarowych umożliwia już ze znacznej odległości rozpoznanie typu danego wehikułu.

Główny korpus (2) magnokraftu czteropędnikowego zajmowany jest przez jego przestrzeń życiową. Korpus ten jest hermetycznie osłonięty poszyciem wykonanym z materiału nieprzenikalnego dla pola magnetycznego (tj. odznaczającego się własnością nazywaną "magnetorefleksyjnością" - patrz opisy z podrozdziału F2.4). Stąd wewnątrz wehikułu czteropędnikowego jest zabezpieczone przed dostępem do niego niebezpiecznego pola magnetycznego. Zawarte w tym korpusie pomieszczenia (jak kabina załogi, pomieszczenia dla pasażerów, przestrzeń bagażowa), ich wyposażenie (np. komputer pokładowy, urządzenia nawigacyjne i pokładowe, system podtrzymywania życia), oraz zapasy; wszystko to osłonięte jest więc przed niszczycielskim działaniem potężnego pola magnetycznego wytwarzanego przez pędniki wehikułu.

Ściany korpusu wehikułu, oraz osłanianej przez ten korpus przestrzeni życiowej, wykonane są z lustro-podobnego, przezroczystego materiału, którego stopień przezroczystości i odbicia światła może zostać regulowany przez załogę. Stąd podczas lotów w nocy, załoga wehikułu może uczynić te ścianki całkowicie przezroczyste, zamieniając swój wehikuł w rodzaj domku ze szkła. Natomiast przy lotach w przestrzeni kosmicznej w pobliżu słońca, czy lotach w atmosferze przy słonecznej pogodzie, załoga może zamieniać ścianki statku w srebrzyste lustra całkowicie odbijające padające na nie światło, tak że we wnętrzu wehikułu będzie wtedy panował przyjemny półcień. W pozostałych przypadkach lotów, ścianki te mogą zostać wysterowane na przyjęcie dowolnego stanu pomiędzy tymi dwoma skrajnościami. Nie istnieje więc potrzeba aby magnokraft czteropędnikowy dodatkowo wyposażać w okna. Niemniej, aby umożliwić załodze i pasażerom wchodzenie na pokład i opuszczanie statku, wehikuł ten musi posiadać drzwi.

W środku podłogi wehikułu czteropędnikowego pierwszej generacji znajduje się kwadratowy lub prostokątny otwór normalnie zamknięty dzwiami o formie kłapy czy zasuw. Otwór ten jest niewidoczny na rysunku D1 ale pokazany na rysunku Q1. Służy on wsunięciu w niego części piramidalnego dachu następnego wehikułu czteropędnikowego w przypadku gdy kilka tych wehikułów sprzęga się razem w odpowiednik latającego cygara (takie cygaro, tyle że uformowane z magnokraftów dyskoidalnych, pokazane zostało na rysunku F7 i w części #1 rysunku F6). Stąd otwór ten umożliwia sprzęganie ze sobą szerego wehikułów czteropędnikowych. Ponadto używany jest on również jako spodnie wejście na pokład tego statku.

Beczko-podobne lub amforo-podobne pędniki (3) magnokraftu czteropędnikowego zajmują wszystkie cztery naroża jego prostopadłościennego lub sześciennego korpusu. Każdy z tych pędników wytwarza swoją własną kolumnę wirującego pola magnetycznego, której rdzeń na rysunku D1 oznaczony został jako (4). Z powodów omówionych w dalszej części tego opracowania (patrz podrozdział D4), kolumny te będą wyraźnie widoczne dla postronnego obserwatora jako rodzaj czarnych, wirujących, ogromnych wiertel.

Ogólna konstrukcja pędnika magnokraftu czteropędnikowego pierwszej generacji pokazana została w części (b) rysunku D1. Pędnik taki zawiera w sobie pięć komór oscylacyjnych posiadających taki sam, kwadratowy przekrój poprzeczny. Komory te zestawione zostają razem w konfigurację krzyżową już omówioną w podrozdziale C7.2. W konfiguracji pierwszej generacji jedna z komór, nazywana komorą główną (patrz M na rysunkach C9 i D1), umieszczona jest w centrum i następnie otoczona przylegającymi do niej

czterema komorami bocznymi (patrz U, V, W i X na rysunkach C9 i D1). Natomiast w konfiguracjach drugiej oraz trzeciej generacji komora główna otoczona będzie odpowiednio ośmioma lub szesnastoma mniejszymi od niej komorami bocznymi. W konfiguracji krzyżowej pierwszej generacji pokazanej na rysunku D1, komora główna jest czterokrotnie dłuższa od każdej z komór bocznych, ponieważ jej objętość zawsze musi być równa ich sumarycznej objętości. Od zewnątrz konfiguracja krzyżowa komór oscylacyjnych z każdego pędnika okrywana jest owiewką aerodynamiczną wykonaną z materiału przenikalnego przez pole magnetyczne. Dla wehikułów czteropędnikowych pierwszej generacji owa owiewka aerodynamiczna może nadawać pędnikowi albo kształt pękatej beczki - patrz "1" w części (c) rysunku D1, albo też kształt krótkiej amfory - patrz "2" w części (b) rysunku D1. Natomiast w wehikułach drugiej i trzeciej generacji stosunek długości owej beczki czy amfory do jej największej średnicy będzie się odpowiednio zwiększał, tj. upodobnią się one do jakby długiej kolumny wielościennej opiętej w środku rodzajem pasa z komór bocznych.

D2. Działanie magnokraftu czteropędnikowego

Działanie magnokraftu czteropędnikowego jest nieco odmienne od działania pozostałych napędów magnetycznych omawianych w rozdziałach F i E, tj. dyskoidalnego magnokraftu oraz napędu osobistego. Z drugiej strony działanie to jest także nieco do nich podobne. W wehikule czteropędnikowym każdy z jego pędników jest bowiem zdolny do samodzielnego lotu i manewrowania. Stąd korpus główny tego wehikułu jest unoszony w przestrzeni jakby przez cztery dołączone do niego ale nawzajem niezależne dyskoidalne magnokrafty, lecące po równoległych trajektoriach. Każdy z pędników tego wehikułu wytwarza też własną kolumnę wirującego pola magnetycznego.

Zestawienie komór oscylacyjnych pędnika wehikułu czteropędnikowego w konfigurację zwaną "konfiguracja krzyżowa" nadaje mu szereg unikalnych cech jakie poprzednio zapewniane były tylko przez układ napędowy całego dyskoidalnego magnokraftu - porównaj podrozdziały C7.2 i F7.2. Przykładowo jeden taki pędnik jest zdolny do samodzielnego produkowania wirującego pola magnetycznego którego wszystkie parametry mogą być ściśle kontrolowane. Stąd nawet jeśli działając w odosobnieniu od pozostałych pędników danego wehikułu, ciągle byłby on zdolny do kontrolowania swego lotu i manewrów. W wielkim uproszczeniu można więc powiedzieć, że latanie magnokraftem czteropędnikowym polega na koordynowaniu pracy wszystkich czterech jego pędników zachowujących się jak niezależne wehikuły, tak aby wynikowy efekt ich działania powodował popychanie korpusu statku w pożądanym kierunku. Niemniej, jak to prawdopodobnie czytelnik uświadomi sobie z treści niniejszego rozdziału, sterowanie wehikułem czteropędnikowym jest wielokrotnie bardziej złożone niż sterowanie dyskoidalnym magnokraftem.

Pędniki wehikułu czteropędnikowego są w stanie wytworzyć dwa rodzaje wirów magnetycznych: własny i wehikułu. **Wir własny** jest bezustannie wytwarzany przez każdy z pędników i polega on na wprowadzaniu pola magnetycznego produkowanego przez ten pędnik w lokalny ruch wirujący następujący wokół jego osi magnetycznej "m". Na rysunku D1 takie cztery wiry własne pędników oznaczone zostały jako kolumny (4) wirującego pola magnetycznego. **Wir wehikułu** włączany jest jedynie w szczególnych przypadkach (np. szybkich lotów na dużych wysokościach lub w próżni kosmicznej) i powstaje on gdy wszystkie cztery pędniki statku kooperują z sobą (tj. pulsują z wzajemnym 90 przesunięciem fazowym) formując wynikowe zawirowanie pola magnetycznego jakie w swym ruchu obiegowym wiruje naokoło korpusu statku. Jednakże zasada formowania owego wynikowego zawirowania pola w wehikule czteropędnikowym jest odmienna niż zasada formowania wiru magnetycznego w dyskoidalnym magnokrafcie. Wytwarza ona bowiem zjawisko wyporu magnetycznego zamiast zjawiska rotowania obwodów magnetycznych. Ponadto wirowanie linii sił pola tego statku następuje wzdłuż odmiennych trajektorii. Stąd wynikowy wir wehikułu czteropędnikowego nie jest tak efektywny jak wir formowany przez dyskoidalny magnokraft. Z ledwością wystarcza

więc on do napędzania wehikułu w kierunkach wschód-zachód, oraz do uformowania pancerza indukcyjnego jaki osłania ten statek przed obiektami materialnymi kierowanymi w niego (pociskami, meteorytami). Jest on jednak niewystarczający dla wytworzenia efektywnego bąbla próżniowego. Z tego też powodu, jak to zostanie podkreślone w dalszej części tego rozdziału, wehikuł czteropędnikowy nie będzie się odznaczał żadnymi z własności jakich powstanie uzależnione jest od zaistnienia bąbla próżniowego.

Wszystkie pędniki magnokraftu czteropędnikowego wytwarzają niezwykle potężne pole magnetyczne. Jednocześnie, bieguny jednoimienne tych pędników zwrócone są w tym samym kierunku (np. biegun N każdego pędnika ku dachowi wehikułu). Stąd, gdyby ich wydatek pozostawał niewirującym, wtedy musiałyby one nawzajem się odpychać z ogromnymi siłami. Jednakże ponieważ ich wydatek wiruje, wytwarzają one relatywistyczne zjawisko opisane poniżej jakie neutralizuje siły ich wzajemnego odpychania się od siebie. W ten sposób, siłowa stabilność magnokraftu czteropędnikowego uzyskiwana jest na drodze dynamicznej (nie zaś statycznej jak w dyskoidalnym magnokrafcie). Podstawowym więc wymogiem wzajemnej neutralizacji odpychania magnetycznego pomiędzy poszczególnymi pędnikami tego statku jest, iż pole magnetyczne produkowane przez każdy z jego pędników musi nieustannie wirować, nawet jeśli cały wehikuł zawisa bez ruchu.

Relatywistyczne zjawisko wykorzystane do neutralizacji oddziaływań pomiędzy pędnikami statku czteropędnikowego jest dosyć dobrze znane osobom obznajomionym z magnetyzmem. Polega ono na powiększaniu się efektywnej długości magnesu w przypadku szybkiego wirowania jego linii sił wokół ich osi magnetycznej - patrz też podrozdział F5.3 (lub G5.3 w [1a]). Jeśli owe linie sił pola magnetycznego wirują wystarczająco szybko wokół osi centralnej takiego magnesu, ich zakrzywienie zaczyna zacieśniać (kurczyć) się dośrodkowo ku tej osi, zaś w wyniku końcowym wydatek magnesu zostaje stopniowo ograniczony do niewielkiego obszaru rozciągającego się wzdłuż tej osi magnesu. To z kolei przemienia początkowo krótki magnes o szeroko rozbiegającym się polu, w magnes którego pole jest bardzo długie ale zawężone w formę cienkiego pręta. Oczywiście nie jest możliwym mechaniczne zawirowanie całych pędników do szybkości wystarczająco wysokich aby wytwarzane przez nie pole zawężyło się do słupów o grubości mniejszej od $0.5 \cdot l$, tj. połowy rozstawu " l_b " lub " l_w " osi magnetycznych omawianego tu statku kosmicznego - patrz rysunek D1. Jednakże konfiguracja krzyżowa z pędników wehikułu czteropędnikowego symuluje takie wirowanie za pośrednictwem formowania rotującej fali magnetycznej, podobnej do fali wytwarzanej przez pędniki boczne dyskoidalnego magnokraftu (patrz opisy w podrozdziałach F7.2 i C7.2). Owa fala wiruje wokół centralnej osi magnetycznej "m" każdego pędnika. Ponieważ zdolna jest ona do osiągnięcia każdej wymaganej prędkości kątowej, jej sterowanie powoduje kontrolowane formowanie opisanego tu zjawiska relatywistycznego jakie utrzymuje wehikuł czteropędnikowy w dynamicznej stabilności siłowej.

D3. Własności magnokraftu czteropędnikowego

Różnice w działaniu wehikułu czteropędnikowego w porównaniu z działaniem dyskoidalnego magnokraftu, powodują powstanie różnic we własnościach obu tych statków. Generalnie rzecz biorąc, wehikuł czteropędnikowy nie jest w stanie wytworzyć efektywnego bąbla próżniowego wokół swej powierzchni (patrz podrozdział F10.1). Stąd wszystkie własności związane z istnieniem tego bąbla nie wystąpią w tym wehikule. Dla przykładu jego loty będą się łączyły z tarciem o atmosferę oraz z efektami dźwiękowymi formowanymi przez takie tarcie (przykładowo z głośnym "**bangiem**" podczas przechodzenia przez barierę dźwięku). Stąd szybkość wehikułu w atmosferze będzie także ograniczana przez barierę cieplną. Jednakże w przestrzeni kosmicznej szybkość tego wehikułu może być zbliżona do szybkości światła. Nieobecność bąbla próżniowego osłaniającego statek będzie też czynić niemożliwymi jego loty poprzez przedmioty stałe (takie jak skały czy budynki). Manewrowość

wehikułu czteropędnikowego będzie na zbliżonym poziomie jak manewrowość dyskoidalnego magnokraftu. Natomiast jego zdolność do jonizowania otaczającego powietrza będzie mniejsza, stąd również jego obraz jonowy będzie posiadał znacznie inny kształt i elementy charakterystyczne. Przykładowo podczas wznoszenia się wehikułu ów obraz będzie zawierał cztery bardzo wyróżniające się kolumny zjonizowanego powietrza formowane przez cztery pędniki statku. Kolumny te dostawione będą do obwodu wynikowego wiru wehikułu, jaki domko-kształtny korpus statku otoczy zaokrągloną chmurą wirującej plazmy (patrz rysunek Q3). Owa domko-kształtna chmura wirującej plazmy będzie mniej intensywna niż cztery zjonizowane kolumny odchodzące od pędników, ponieważ natężenie wytwarzającego ją pola jest też mniejsze. Podczas opadania wehikułu czteropędnikowego, słupy zjonizowanego powietrza wytwarzane przez wiry własne jego pędników mogą zaniknąć, stąd jedynie wynikowa, chatko-kształtna chmura otaczająca cały wehikuł może pozostać widoczna.

Kilka wehikułów czteropędnikowych jest w stanie łączyć się z innymi statkami magnetycznymi, formując w ten sposób wiele różnorodnych konfiguracji latających znanych już z opisów dyskoidalnego magnokraftu. Przykładowo para (dwa) lub więcej tych wehikułów może połączyć się razem formując odpowiednik cygaro-kształtnego kompleksu latającego (patrz część #1 rysunku F6), lub odpowiednik kompleksu kulistego (patrz rysunek F1c). Również grupa cygar uformowanych w ten sposób może łączyć się dalej w którąś z konfiguracji wyższego rzędu, reprezentujących odpowiednik dla latającego systemu lub latającego klustera dyskoidalnych magnokraftów (patrz części #5 i #6 rysunku F6).

Wehikuły czteropędnikowe mogą również łączyć się z dyskoidalnymi magnokraftami w różnorodne konfiguracje latające. W połączeniach takich przywierają one do tych statków w sposób zapewniający, iż wyloty z ich pędników ustawiają się precyzyjnie na wylotach z bocznych pędników dyskoidalnych magnokraftów. Aby umożliwić takie ustawienie, wehikuły czteropędnikowe będą budowane jedynie we wielkościach jakie odpowiadają wymiarom poszczególnych typów dyskoidalnych magnokraftów (tj. jakie umożliwiają ułożenie się osi pędników danego wehikułu z osiami pędników odpowiadającego mu typu dyskoidalnego magnokraftu). Z tego też powodu również opracowanych zostanie tylko osiem podstawowych typów wehikułu czteropędnikowego. Wymiary dla tych typów zestawiono w **tablicy D1**. Poszczególne typy tego wehikułu oznakowane zostały jako T3, ..., do T10. Każdy z nich posiada rozstaw osi magnetycznych swych pędników rozłożony dokładnie wzdłuż obwodu okręgu o średnicy nominalnej "d" jaka precyzyjnie pokrywa się z rozstawem pędników bocznych u takiego samego typu dyskoidalnego magnokraftu. Przykładowo typ T3 wehikułu czteropędnikowego posiada swe pędniki ustawione dokładnie na wylotach pędników bocznych w typie K3 dyskoidalnego magnokraftu, typ T4 - w typie K4, itd.

Podobnie jak dyskoidalny magnokraft, również magnokraft czteropędnikowy podczas lądowania wypali w glebie charakterystyczne ślady zwane "lądowiskami". Lądowiska te składać się będą z czterech kolistych wypaleń na ziemi, spowodowanych przez każdy z czterech pędników tego wehikułu - patrz (6) na rysunku D1. Rozłożenie tych wypaleń odpowiadać powinno w przybliżeniu narożnikom czworoboku o wymiarach nieco większych od wymiarów danego statku (tj. wymiary statku powinny dać się wpisać w obręb uformowanych przez niego śladów). Jednakże z uwagi na fakt że wehikuł czteropędnikowy zwykle zawisa nieco nachylony, a także że osie magnetyczne jego poszczególnych pędników nie są dokładnie równoległe do siebie (np. aby kompensować reakcyjny moment obrotowy od wiru wehikułu), wzajemne położenie śladów wypalonych w glebie może znacznie odbiegać od kształtu idealnego kwadratu lub prostokąta. Stąd w rzeczywistości ślady pozostawione przez wehikuł czteropędnikowy zwykle przyjmowały będą postać czterech podobnych wypaleń formujących narożniki czworokąta nierównobocznego (tj. czworokąta którego każdy z boków posiada inną długość formując obrys zdeformowanego prostokąta). Jeśli chodzi o kształt poszczególnych wypaleń pozostawianych przez cztery pędniki tego wehikułu, to zależał on będzie od trybu pracy użytych w nich konfiguracji krzyżowych. Przy ich trybie pracy z "dominacją strumienia wewnętrznego" (patrz opis w podrozdziale C7.2), wypalane zostaną w glebie charakterystyczne ślady składające się ze silnie zaznaczonego centralnego wypalenia i

mniej wyraźnej pierścieniowatej obwódki. Takie właśnie ślady zilustrowano na rysunku D1 (wskazano je tam odnośnikiem "6"). W przypadku trybu pracy pędników z "dominacją strumienia zewnętrznego", każdy z nich uformuje ślad w kształcie silnie wypalonego pierścienia (obwódki) z mniej wypalonym centralnym obszarem. Warto tu też dodać, że teoretycznie rzecz biorąc część pędników (np. dwa) wehikułu czteropędnikowego mogłaby pracować w dominacji strumienia wewnętrznego, część zaś (np. pozostałe dwa) w dominacji strumienia zewnętrznego. W takim przypadku wypalenia formowane przez nie w glebie mogłyby przyjąć postać mieszaniny obu powyżej omówionych śladów. Niemniej w praktyce tak duża elastyczność działania pędników jest prawdopodobnie niemożliwa z uwagi na niepomiarnie większą złożoność komputerowych programów sterujących danym wehikułem wymaganą dla jej realizacji. Dla zapewnienia tej elastyczności programy owe musiałyby być bowiem wielokrotnie obszerniejsze i bardziej uniwersalne od programów sterujących zezwalających wszystkim pędnikom jedynie na pracę w tym samym trybie (stąd też, dzięki swojej złożoności, mogłyby one ukrywać znacznie więcej błędów wprowadzając poprzez to również zwiększone niebezpieczeństwo katastrofy sterowanego nimi wehikułu - patrz [5]). Stąd też najprawdopodobniej takie mieszane ślady nie będą formowane w praktyce.

D4. Wygląd i wymiary magnokraftu czteropędnikowego

Wygląd magnokraftu czteropędnikowego pokazany został w części (a) rysunku D1. Z uwagi na dynamiczne neutralizowanie w nim międzypędnikowych oddziaływań magnetycznych, kształt i wymiary przestrzeni życiowej wehikułu czteropędnikowego nie są ograniczane warunkami stabilności (jak to było w przypadku dyskoidalnego magnokraftu - patrz podrozdział F4.3). Stąd, głównymi kryteriami konstrukcyjnymi przy projektowaniu kształtu tego wehikułu stają się: (1) zagwarantowanie magnetycznego sprzęgania się tych wehikułów z dyskoidalnymi magnokraftami w konfiguracji latające; (2) zapewnienie możliwie najszybszej i najłatwiejszej identyfikacji typu danego wehikułu czteropędnikowego; (3) zapewnienie wymaganej aerodynamiczności lotów w ośrodkach gazowych i płynnych; (4) dostarczenie możliwie najwyższego komfortu lotu dla załogi i pasażerów; (5) zabezpieczenie najłatwiejszego lądowania; (6) wspomaganie możliwie najłatwiejszego załadunku, przenoszenia i rozładunku transportowanego bagażu, itp.

Powyższe kryteria konstrukcyjne zezwalają na określoną swobodę w projektowaniu kształtu korpusu wehikułu czteropędnikowego. Z uwagi jednak na najwyższą przydatność kształtu sześciennego do transportowania ludzi i bagażu, wehikuł ten najczęściej będzie przyjmował właśnie kształt kostki sześcienniej (po angielsku "cubicle") z daszkiem przypominającym małą piramidkę - taki jak pokazany na rysunku D1 (a). W niektórych przypadkach może on też przyjąć formę prostopadłościennego domku (chatki) również z piramidkowym dachem - taki jak pokazany na rysunku Q1. Oczywiście w szczególnych przypadkach dowolne inne kształty mogą też zostać użyte, dla przykładu kuliste, rakietopodobne, stożkowe czy nawet aerodynamiczne jak u współczesnego samolotu (ponieważ te odmienne kształty będą używane raczej rzadko, omówienie i ilustracja ich charakterystyk użytkowych zostały pominięte w niniejszym rozdziale).

Kryteria numer (1) i (2) wymienione powyżej narzucają szereg warunków matematycznych na wymiary wehikułu czteropędnikowego. Aby usatysfakcjonować owe warunki wymiary te muszą spełniać różne współzależności i równania jakich wyprowadzenie zostanie tu pominięte, jakie jednak przytoczone zostały u spodu tablicy D1. Z uwagi na owe współzależności, wymiary poszczególnych typów wehikułów czteropędnikowych będą więc ściśle zdefiniowane. Ich wartości zestawione zostały w tablicy D1. Posiadanie owej tablicy umożliwia aby w przypadku ustalenia wartości jakiegokolwiek jednej wielkości dotyczącej danego wehikułu czteropędnikowego (np. liczby członków jego załogi), natychmiast stało się też możliwe odczytanie pozostałych danych opisujących ten wehikuł, takich jak jego wymiary, pędniki, waga, itp.

Kryterium konstrukcyjne numer (1), dotyczące możliwości sprzęgania się wehikułów czteropędnikowych z dyskoidalnymi magnokraftami, decyduje także o kształcie i wymiarach pędników użytych w wehikułach czteropędnikowych. Powodem jest tutaj wymóg, aby po sprzęgnięciu się z dyskoidalnym magnokraftem odpowiadającego typu, gęstość energii w pędnikach pozostawała taka sama. To z kolei narzuca wymóg, aby suma objętości (V_c) wszystkich komór oscylacyjnych z pędników wehikułu czteropędnikowego była równa sumie objętości (V_d) wszystkich komór oscylacyjnych z pędników dyskoidalnego magnokraftu odpowiadającego mu typu, tj. aby: $V_c = V_d$. To z kolei narzuca aby wymiary boczne komór oscylacyjnych obu tych statków miały się do siebie w odpowiedniej proporcji, jaką stosunkowo łatwo daje się wyliczyć. Przykładowo jeśli przez symbol " a_M " oznaczyć wymiar boku zewnętrznej komory oscylacyjnej z pędnika głównego (M) dyskoidalnego magnokraftu (zestawiony w kolumnie " a_M " tabeli F1), zaś przez " a " wymiar boku w komorach oscylacyjnych standardowej konfiguracji krzyżowej użytej w wehikułach czteropędnikowych (patrz rysunek D1), wtedy wzajemna zależność pomiędzy tymi wymiarami opisana może zostać następującym przybliżonym wzorem wynikającym z warunku $V_c = V_d$:

$$a = 0.397 a_M \quad (D1)$$

(Zależność tą daje się też wyrazić następującym dokładnym wzorem (D1'): $a = (1/16)^{1/3} a_M$.)

Istotnym elementem wyglądu wehikułu czteropędnikowego są długie, cienkie kolumny wirującego pola magnetycznego produkowanego przez każdy z jego pędników. Ponieważ kolumny te posiadają wyraźnie zaznaczające się granice, zaś formujące je pole o ogromnej koncentracji jest zawsze szybko-pulsujące, stąd będą one stanowiły rodzaj pułapki dla światła (patrz opis "czarnych belek" z podrozdziału F10.4 tej monografii i G3.4 monografii [1a]). Podczas dnia dla przypadkowego obserwatora będą więc one wyglądały jakby zostały wykonane z czarnego materiału lub dymu. Ich ciągły ruch wirowy na przypadkowym obserwatorze może sprawiać wrażenie oglądania układu czterech czarnych wiertel wirujących z ogromnymi szybkościami.

Zupełnie odmienny wygląd te kolumny pola przyjmą jeśli będą one oglądane w nocy. Ponieważ jonizują one powietrze, ich wygląd na tle czarnego otoczenia będzie wtedy przypominał tzw. "biały szum" oglądany na ekranie telewizora (biały szum - po angielsku "white noise" jest to obraz złożony z białych i czarnych kropek jaki ukazuje się na ekranie telewizora jeśli telewizor ten pozostaje włączony ale nie odbiera żadnej stacji).

W każdej kolumnie pola odprowadzanego do otoczenia z pędników wehikułu czteropędnikowego daje się wyróżnić dwa obszary, tj.: ciemny rdzeń (4) i jaśniejszą otoczkę (5) w przypadku pracy tych pędników w trybie dominacji strumienia wewnętrznego - jak to pokazano w części (a) na rysunku D1; lub jaśniejszy rdzeń i ciemną otoczkę w przypadku pracy tych pędników w trybie dominacji strumienia zewnętrznego. Rdzeń (4) uformowany zostaje w efekcie wirowania wydatku z głównej komory oscylacyjnej w konfiguracji krzyżowej danego pędnika wokół osi magnetycznej "m" tej komory (na rysunku D1 komora główna oznaczona jest jako M). Natomiast szybko obracające się cztery ramiona otoczki (5) formowane są w efekcie wirowania wokół rdzenia (4) wydatków z czterech bocznych komór oscylacyjnych danej konfiguracji krzyżowej (na rysunkach D1 i C9 oznaczonych jako U, V, W i X). Wydatki tych komór bocznych przylegają do rdzenia i wirują wraz z nim w podobny sposób jak pióra wiertła przylegają i wirują wraz z rdzeniem tego wiertła.

Jak to już wyjaśniono powyżej, podczas dnia dla zewnętrznego obserwatora wygląd owych dwóch części wirującej kolumny pola pędnika czyni to pole niezwykle podobne do wirującego czarnego wiertła. Owo wiertło z kolei nie jest zbyt odległe w wyglądzie od czarnego śmigła helikopterowego, tyle tylko iż zamiast być płaskim i szerokim jest ono wąskie i długie. To, razem z kanciastym, helikoptero-podobnym kształtem samego wehikułu, jest w stanie spowodować u niektórych - mniej obznajomionych z techniką lotniczą widzów, iż niekiedy nawet bez użycia modyfikatorów wyglądu opisanych w punkcie 5 podrozdziału N3.2, ciągle mogą oni pomylić wehikuł czteropędnikowy z dwu- lub cztero-rotorowym helikopterem pozbawionym znaków rozpoznawczych (symboli rejestracyjnych).

D5. Identyfikacja typu magnokraftu czteropędnikowego

Aby umożliwić innym podróżnikom przestrzeni kosmicznej szybkie identyfikowanie typu napotkanego przez nich magnokraftu czteropędnikowego, najważniejsze elementy geometryczne tego wehikułu budowane będą w odpowiednich proporcjach. W ten sposób identyfikowanie typu danego wehikułu stanie się niezwykle łatwe i stąd może być dokonywane nawet przez komputer podłączony do radaru. Identyfikowanie to wymaga jedynie wyznaczenia wzajemnej proporcji pomiędzy najważniejszymi wymiarami wehikułu. Z kolei, znajomość tych proporcji wskaże współczynnik typu "T" danego wehikułu czteropędnikowego. Wartość owego współczynnika typu jest równa stosunkowi gabarytowej wysokości "H" danego wehikułu do wysokości "Z" jego piramidkowego dachu (patrz rysunek D1):

$$T = H/Z$$

(D2)

Kiedy z proporcji wymiarowych danego wehikułu czteropędnikowego poznany zostanie jego współczynnik T, wszystkie pozostałe dane o tym wehikule mogą albo zostać odczytane z odpowiednich tablic (patrz tablica D1), albo też wyznaczone z odpowiednich współzależności matematycznych obowiązujących dla tych statków (patrz wzory zestawione pod tablicą D1).

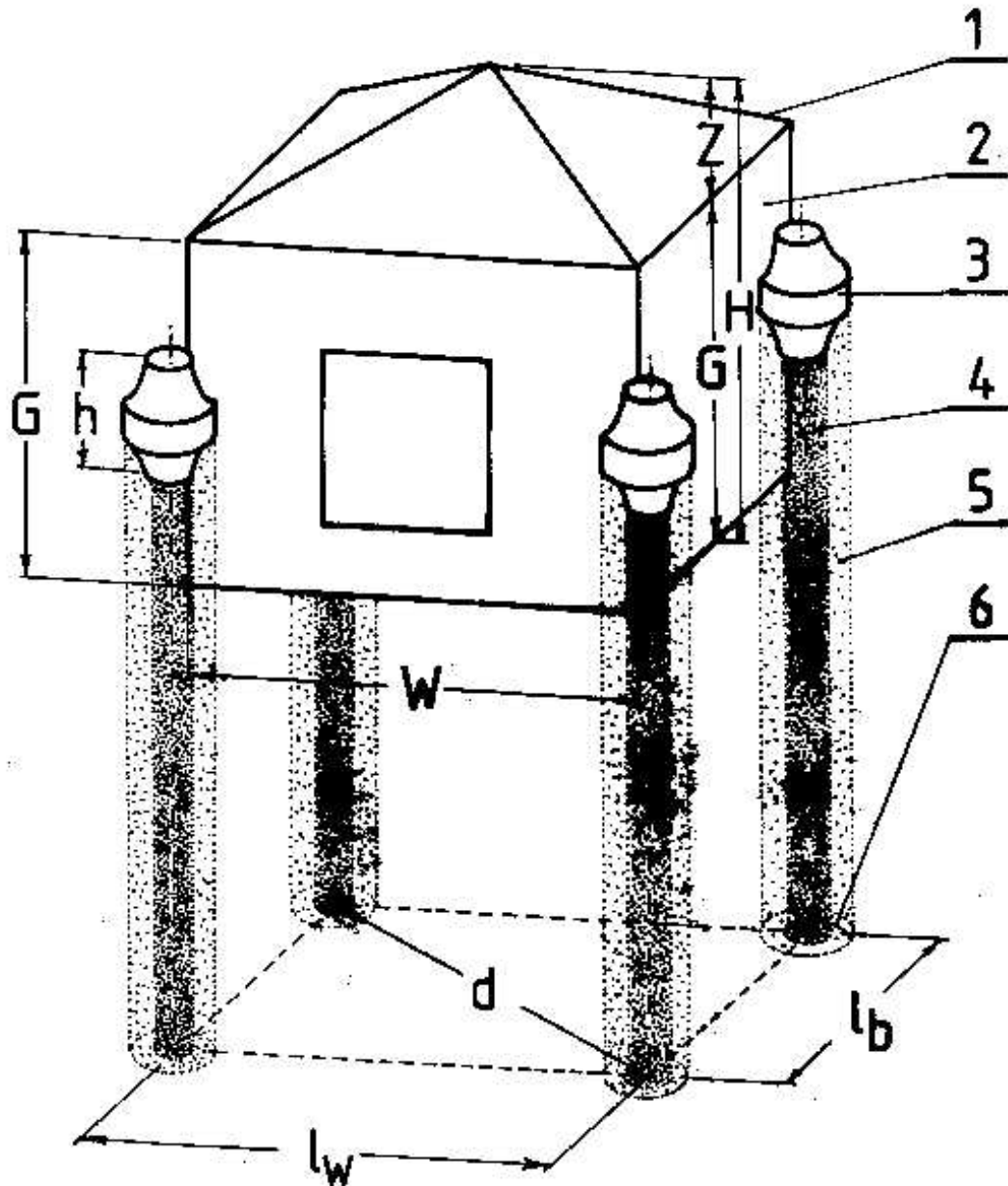
D-103

Tablica D1. **Najważniejsze dane konstrukcyjne wszystkich ośmiu typów magnokraftów czteropędnikowych.** Interpretacja użytych symboli pokazana została na rysunku D1. Wymiary poszczególnych wehikułów wyznaczono z warunku ich sprzęgalności z dyskoidalnymi magnokraftami, (tj. w wehikułach sześciennych dla których $l=l_b=l_w$, odległość "l" pomiędzy osiami ich pędników musi spełniać następujące równanie: $l = 0.5486 \cdot 2^{(T-1)}$ [metrów]). Wszystkie wymiary liniowe z tej tablicy wyrażone zostały w metrach.

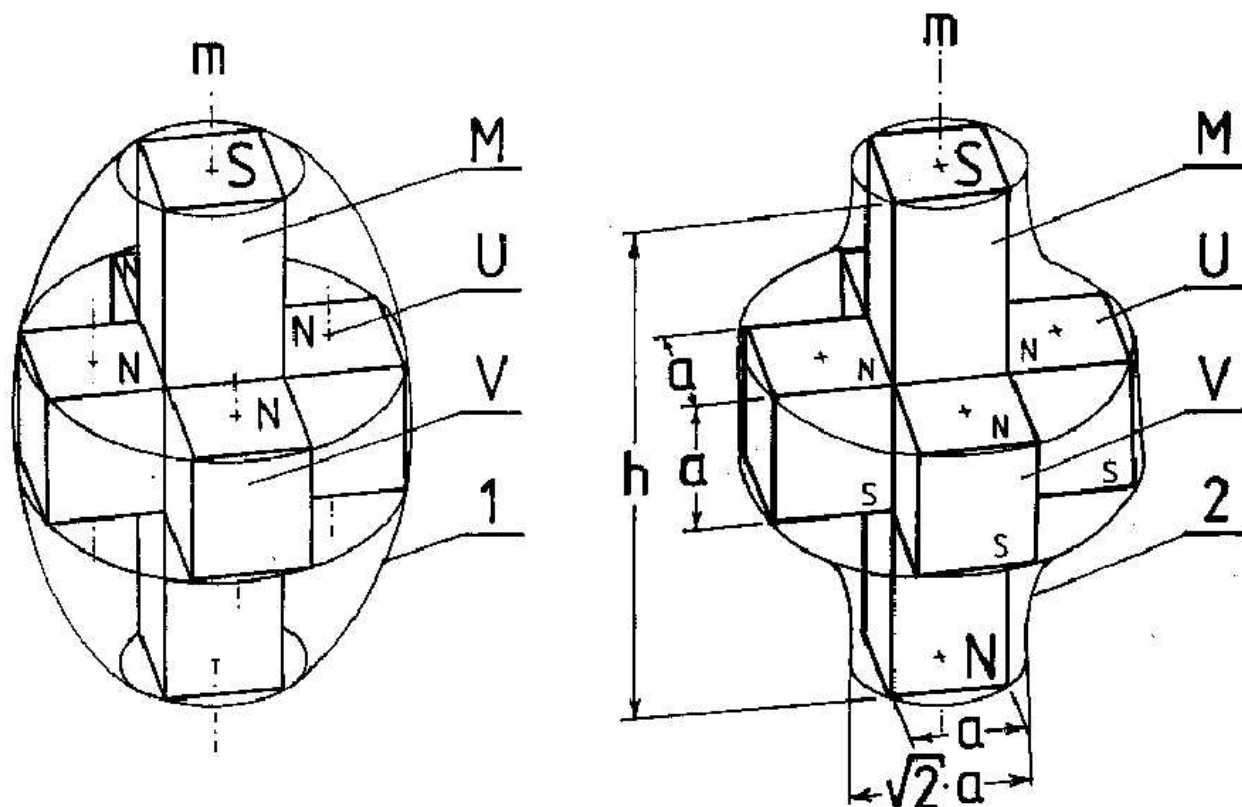
Nr	Typ	Odpow typ dysk.	Wymiary gabarytowe korpusu wehikułu dla statków sześciennych				Rozstawienie pędników wehikułu				Wymiary pędników		Za ło ga	Waga stat- ku	
							d (prze- kątna)	Statki prostokątne		sześć. l=lw,b					
	T	K	W	G	Z	H		ANG	lw		lb	h			a
-	-	-	m	m	m	m	m	deg	m	m	m	m	m	-	ton
1.	T3	K3	2.01	1.46	0.73	2.19	3.10	22.5	2.86	1.19	2.19	0.73	0.18	3	0.5
2.	T4	K4	4.11	3.29	1.09	4.38	6.20	30	5.37	3.10	4.39	1.09	0.27	4	4
3.	T5	K5	8.35	7.02	1.76	8.78	12.41	33.75	10.32	6.89	8.78	1.76	0.43	5	33
4.	T6	K6	16.82	14.64	2.93	17.55	24.82	27	15.64	7.97	17.56	2.93	0.73	6	270
5.	T7	K7	33.86	30.09	5.02	35.11	49.65	30	43.00	24.82	35.11	5.02	1.25	7	2 164
6.	T8	K8	68.02	61.44	8.78	70.22	99.30	32.14	59.46	37.36	70.22	8.78	2.20	8	17 312
7.	T9	K9	136.54	124.84	15.60	140.44	198.61	28.125	123.86	66.20	140.44	15.60	3.90	9	138 497
8.	T10	K10	273.86	252.79	28.09	280.88	397.22	30	344.00	198.61	280.88	28.09	7.02	10	1107981

Oto równania wyrażające związki matematyczne pomiędzy poszczególnymi wielkościami z tej tablicy:

$$\begin{array}{llllllll}
 T=H/Z & T=K & Z=h & d=l\sqrt{2} & a=h/4 & d^2 = l_w^2 + l_b^2 = 2l^2 & l=0.5486 \cdot 2^{(T-1)} \text{ [m]} \\
 Z=H/T & H=l & Z=l/T & ANG=\arctan(l_b/l_w) & h=l/T & Waga=0.05 \cdot l^2 \cdot H
 \end{array}$$



Rys. D1. Wygląd ogólny wehikułu czteropędnikowego. Statek ten, razem z dyskoidalnym magnokraftem oraz magnetycznym napędem osobistym, reprezentuje jedno z trzech podstawowych zastosowań pędników magnetycznych wykorzystujących komorę oscylacyjną. Zilustrowane zostały: kształt, podzespoły, oraz najważniejsze wymiary tego wehikułu. Symbole: 1 - dach w kształcie piramidki; 2 - sześcienny korpus główny statku zawierający jego przestrzeń życiową (tj. kabinę załogi, kabiny pasażerów, powietrze, zapasy, komputer pokładowy, itp.); 3 - jeden z czterech pędników; 4 - rdzeń słupa pola magnetycznego wydzielanego przez każdy z pędników tego wehikułu (rdzeń ten formowany jest z pola produkowanego przez główną "M" komorę oscylacyjną); 5 - otoczka z wirujących segmentów pola magnetycznego wydzielanego z komór bocznych U, V, W, X każdego pędnika; 6 - jeden z czterech wypalonych śladów pozostawianych na powierzchni gruntu przez taki nisko zawieszający statek którego pędniki pracują w trybie dominacji strumienia wewnętrznego (patrz podrozdziały D3 i C7.2). Wymiary: H, Z, G, W - opisują rozmiary prostokątnej lub sześcienniej kabiny załogi (reprezentują one: wysokość gabarytową, wysokość dachu, wysokość ścian, oraz szerokość statku); d, l, l_w , l_b (dla sześciianu $l_w = l_b = l$) - opisują rozstaw osi magnetycznych wehikułu (rozstaw ten musi być zgodny z rozstawem pędników bocznych dyskoidalnego magnokraftu tego samego typu); h - opisuje wysokość pędników wehikułu.



Rys. D2. Pędniki wykorzystywane w wehikule czteropędnikowym. Rysunek pokazuje kształt, wygląd, oraz główne wymiary owych pędników. Symbole: M - komora główna danej konfiguracji krzyżowej; U, V, W, X - cztery komory boczne konfiguracji krzyżowej pierwszej generacji przyjmujące kształty sześcianów (zauważ że przekrój poprzeczny komór U, V, W, X musi się równać przekrojowi poprzecznemu komory M, jednak ich długość musi być czterokrotnie mniejsza); 2 - amforo-kształtna owiewka aerodynamiczna jaka przykrywa i osłania komory pędnika (zauważ że owiewka ta może też przyjąć inne kształty); a - wymiar boczny sześciennej komory oscylacyjnej; N, S - zorientowanie biegunów magnetycznych w poszczególnych komorach oscylacyjnych; m - oś magnetyczna pędnika.

(**Lewy**) Pędnik beczko-kształtny pierwszej generacji. Jedyna różnica w stosunku do pędnika z prawej części, to inny kształt przenikalnej dla pola magnetycznego owiewki aerodynamicznej "1".

(**Prawy**) Pędnik amforo-kształtny pierwszej generacji. Pokazano jego wygląd, wymiary, oraz podzespoły składowe. Oparty jest on na konfiguracji komór oscylacyjnych nazywanej "konfiguracja krzyżowa" (po szczegóły patrz rysunek C9).

MAGNETYCZNY NAPĘD OSOBISTY

Motto: "Użyteczność wzrasta proporcjonalnie do stopnia zminiaturyzowania (tj. im mniejsze tym użyteczniejsze)."

Prawie każdy z nas śnił kiedyś o lataniu. Przypomnijmy więc sobie jak ono wyglądało. Nasz mózg nabrzmiewał decyzją wzniesienia się w powietrze, zaś ciało posłusznie i bezwzględnie podążało za zleceniem woli. Nie musieliśmy wymachiwać rękami czy przebiegać nogami. Wszystko co pomyśleliśmy zostało natychmiast wykonane.

Ciekawe skąd się bierze owa rozbieżność naszych snów z logiką która przecież podpowiada, iż latanie zawsze powinno wymagać wymachiwań (wszakże ptaki machaniem skrzydłami ciężko zapracowują na swoje loty). Uważam, że nasza intuicja już obecnie wie co nadejdzie w dalekiej przyszłości. Bezwysiłkowe sny naszego latania prawdopodobnie są więc intuicyjnymi obrazami urządzeń jakich zbudowanie ma dopiero nastąpić.

Jak nasi odlegli potomkowie odbywać będą latanie w powietrzu już obecnie daje się wydedukować na podstawie działania wehikułu opisanego w rozdziale F pod nazwą "dyskoidalnego magnokraftu". Streszczając to działanie w wielkim uproszczeniu, urządzenia napędowe magnokraftu przyjmują formę potężnych źródeł pola (tj. "komór oscylacyjnych") zamkniętych w kulistych obudowach i w tej monografii nazywanych "pędnikami magnetycznymi". Podczas lotów w pozycji odwróconej, pole wytwarzane przez osiem pędników bocznych magnokraftu odpycha ten wehikuł od pola magnetycznego Ziemi, wynosząc go w przestrzeń - patrz pędniki boczne oznaczone U, V, W, X na rysunkach F1 "a", F3, i E1 "a". Równocześnie pojedynczy pędnik główny umieszczony w centrum statku przyciągany jest przez pole ziemskie formując siłę stabilizacyjną jaka utwierdza zorientowanie statku w przestrzeni oraz kontroluje jego wzlot, zawisanie, lub opadanie - patrz pędnik główny oznaczony M na rysunkach F1 "a", F3 i E1 "a". Omawianą w powyższym streszczeniu pozycję odwróconą magnokraftu można by nazwać pozycją "wiszącą" aby ją odróżnić od normalnej pozycji lotów magnokraftu nazywanej "stojącą" i pokazanej na rysunku F4. Magnokraft lecący w owej odwróconej pozycji "wiszącej" pokazano właśnie w części "a" rysunku E1.

Łatwo przewidzieć, że pędniki magnetyczne wykorzystujące komorę oscylacyjną, pewnego dnia zostaną zminiaturyzowane do kilkunastomilimetrowych rozmiarów i następnie zabudowane do elementów odpowiednio adaptowanej garderoby, np. do ośmiosegmentowego pasa i grubych podeszw butów. Jak to widać w części "b" rysunku E1, po zabudowaniu do pasa i butów pędniki te uformują zarys sylwetki ludzkiej a jednocześnie ich działanie będzie niemal identyczne jak działanie pędników w magnokrafcie pokazanym w części "a" tego samego rysunku E1. Dzięki temu takie ich zestawienie pozwoli na uformowanie nowego rodzaju napędu, jaki nazywam "magnetycznym napędem osobistym". Urządzenie wykorzystujące ów nowy napęd znajdzie zastosowanie dla powodowania lotów osób bez użycia żadnego widocznego wehikułu, lub do wspomagania tradycyjnych sposobów poruszania się tych osób (np. do chodzenia po powierzchni wody lub po suficie, do wskakiwania na dachy najwyższych budynków, itp.)

Podobnie jak to jest z pędnikami najmniejszego magnokraftu typu K3, napęd osobisty wykorzystywał będzie osiem pędników bocznych (oznaczonych U, V, W i X) - patrz część "b" rysunku E1. Jednakże w przeciwieństwie do magnokraftu posiadał on będzie nie jeden, ale aż dwa pędniki główne (w części "b" rysunku E1 oznaczone M_L i M_R). Obie te grupy pędników

zamocowane zostaną do ciała użytkownika za pośrednictwem jednoczęściowego kombinezonu, tworząc w ten sposób wysoce efektywny system napędowy. Ciało użytkownika wypełni w nim funkcję "konstrukcji nośnej" lub "ramy". Każdy pędnik takiego systemu, podobnie jak pędnik magnokraftu, zawiera w sobie jedną zminiaturyzowaną kapsułę dwukomorową wielkości zaledwie kilkunastu milimetrów, jaka zamontowana zostaje we wnętrzu odpowiedniej kulistej obudowy. Owa kapsuła oraz jej obudowa są podobnej konstrukcji i działania jak te użyte w magnokrafcie, tyle tylko że zostały one odpowiednio zminiaturyzowane. Dlatego też pędniki napędu osobistego mogą być budowane w wielkościach pozwalających na ich zamontowywanie do wnętrza części garderoby (np. butów i pasa) bez powodowania zauważalnego zwiększenia niewygody, czy ciężaru i wielkości tej garderoby. Z drugiej strony, pozostając prawie że niezauważalnymi, pędniki te dostarczą ich użytkownikom różnorodnych atrybutów wyszczególnionych w podrozdziale E6, takich jak przykładowo: zdolność do latania w powietrzu lub przestrzeni kosmicznej z prędkościami limitowanymi jedynie wykonywaniem czynności fizjologicznych (przykładowo oddychaniem), ogromna siła fizyczna, niewidzialność, odporność na działanie broni palnej i każdej innej broni jaka mogłaby być przeciwko nim użyta, plus wiele innych równie pożądaných i niezwykłych możliwości.

Jeśli chodzi o kolejność budowy poszczególnych napędów magnetycznych na Ziemi, to magnetyczny napęd osobisty zbudowany zostanie jako czwarty rodzaj napędu wykorzystującego komory oscylacyjne (patrz okres 1D z klasyfikacji podanej w podrozdziale M6). Przyczyną tego stanu rzeczy będą początkowe trudności technologiczne z miniaturyzacją kapsuł dwukomorowych do rozmiarów na tyle niewielkich jakie wymagane będą dla tego napędu.

E1. Standardowy kombinezon napędu osobistego

Wygląd standardowego kombinezonu napędu osobistego pokazany został na **rysunku E2**. Wszystkie części tego kombinezonu przypominają typowe elementy ubioru ludzkiego, tyle tylko że niezależnie od funkcji ubiorczych dodatkowo dostosowane są one do wypełniania funkcji transportowych. Dla niezorientowanego obserwatora nie będzie więc możliwym wzrokowe wykrycie istnienia takiego napędu, jako że jego wygląd przypominał będzie zwykły strój sportowy, zaś jego istnienie stanie się zauważalne dopiero po ujawnieniu się wywołanych przez niego efektów napędowych (np. wzniesieniu się jego użytkownika w powietrze). Na kombinezon napędu osobistego składają się następujące elementy: jednoczęściowy kostium (3) z kapturem osłaniającym głowę (5) oraz rękawicami (4), buty (1) zawierające miniaturowe "pędniki główne" zamontowane w ich podeszwach, oraz specjalny ośmiosegmentowy pas (2) utrzymujący wbudowane w niego osiem miniaturowych "pędników bocznych". Kaptur (5), rękawice (4) i buty (1) są tak zaprojektowane aby hermetycznie łączyły się one z kostiumem (3), formując w ten sposób jednoczęściowy kombinezon szczelnie osłaniający całe ciało użytkownika. Z tyłu kołnierza tego kombinezonu wbudowany zostanie komputer sterujący napędem oraz czujniki jakie odczytywać będą myślowe sygnały sterownicze bezpośrednio z podstawy czaszki użytkownika i następnie przetwarzać je na rozkazy wykonawcze w sposób już opisany dla urządzenia zwanego TRI (po szczegóły patrz podrozdział N3).

Wszystkie elementy ubiorcze omawianego kombinezonu wykonywane będą ze specjalnego materiału magnetorefleksyjnego - patrz jego opisy w podrozdziale F2.4. Materiał taki odbija pole magnetyczne w sposób podobny jak lustro odbija światło. Dzięki tej własności, materiał magnetorefleksyjny zabezpieczy ciało użytkownika przed niszczącym działaniem silnego, pulsującego pola magnetycznego wytwarzanego przez pędniki napędu osobistego. Aczkolwiek twarz pozostanie odsłonięta, pole magnetyczne również nie będzie w stanie przez nią wnikać do mózgu, ponieważ działanie tego pola ogranicza się tylko do obszarów gdzie jego linie sił są w stanie formować obwody zamknięte (kształt magnetorefleksyjnego kaptura czyni niemożliwym takie domykanie się obwodów poprzez głowę użytkownika). Dla

dotatkowego zabezpieczenia skóry twarzy przed silnym polem magnetycznym niekiedy koniecznym się stanie dodatkowe ubieranie specjalnej maski, podobnej do masek zakładanych przez rabusiów bankowych lub przez "spiderman", "batman", i "superheroes" z filmów amerykańskich. W przypadku gdy maska taka okazałaby się niestosowna, możliwe jest też użycie kremu wykonanego na bazie grafitu (stwierdzono, że grafit jest jednym z najlepszych naturalnych materiałów magnetoreflekcyjnych). Oczywiście ów krem nie tylko że będzie zabezpieczał twarz przed polem magnetycznym, ale także nada skórze użytkownika niezwykle, metalicznego kolorytu. Jedyne części twarzy które bez użycia hełmu kosmicznego jak ten z rysunku E4 "b" nie będą mogły zostać zabezpieczone przed działaniem na nie pola magnetycznego to oczy i zęby. Stąd oczy i zęby użytkowników tego typu napędu pobudzone silnym polem magnetycznym będą w ciemności świeciły fosforycznym światłem, stwarzając podstawy do powstania różnych opowieści o wampirach i innych opisywanych w podrozdziale R4.1 nadprzyrodzonych istotach.

Specjalne rękawice (4) uzupełniające magnetyczny napęd osobisty zostały tak zaprojektowane aby nie tylko zabezpieczać palce przed wpływem silnego pola magnetycznego, ale także przed indukowanymi przez to pole siłami elektrostatycznymi. Jak można się bowiem spodziewać, pulsujące pole magnetyczne pędników zawartych w pasie musi powodować wytwarzanie silnych ładunków elektrycznych wokół bioder użytkownika. Z kolei ładunki te stopniowo zgromadzą się na palcach użytkownika. Siły wzajemnego odpychania się tych jednoimiennych ładunków będą usiłowały rozchyłać palce użytkownika na podobnej zasadzie jak czynią to z listkami elektroskopu. Aczkolwiek owe rozchylające działanie ładunków będzie zbyt słabe aby zaszkodzić użytkownikowi, przy długotrwałym jego działaniu może ono jednak spowodować nieprzyjemne naciąganie skóry i ból mięśni. Stąd użycie błono-podobnych łączników pomiędzy palcami rękawic (4), podobnych do błon na nogach kaczki, zabezpieczy użytkownika przed owymi nieprzyjemnymi następstwami.

Aby zmniejszyć ilość elektrostatycznej elektryczności gromadzącej się na palcach, użytkownicy napędu osobistego z pędnikami w pasie będą wykazywali też tendencję do trzymania swych rąk w podniesieniu, albo unosząc je z boku jak zrywający się do lotu ptak wznosi swe skrzydła, albo też trzymając je złożone razem i wyciągnięte do przodu jak przypisuje się to lunatykom (patrz też opisy z podrozdziału R4.1). Bardzo pomocne w takim utrzymywaniu rąk w podniesieniu okażą się bransoletki wspomagające opisane w następnym paragrafie i pokazane jako 3 na rysunku E4 "a". Bransoletki te same unosiły będą ręce swych użytkowników, tak że stałe utrzymywanie tych rąk w omówionym tu podniesieniu nie będzie wymagało żadnego wysiłku.

Jeśli użytkownik napędu osobistego zamierzał będzie wykonać ciężką pracę fizyczną, wykorzysta on dodatkowe pędniki wspomagające. Pędniki te zamontowane w specjalnych bransoletkach wspomagających zakładane będą na oba nadgarstki jak zegarki ręczne (patrz 3 na rysunku E4 "a"). Poprzez magnetyczne oddziaływania siłowe z innymi pędnikami napędu osobistego, pędniki wspomagające uformują rodzaj bezdotykowego dźwigu nadającego użytkownikowi niezwyklej siły fizycznej. Stąd osoba w nie wyposażona zdolna będzie do podnoszenia głazów ważących wiele ton, do przełamywania potężnych konstrukcji, powalania budynków, wrywania drzew z korzeniami, itp. Bez żadnego też wysiłku będzie godzinami mogła utrzymywać swe ręce w podniesieniu jak to opisano jeden paragraf wyżej.

W niektórych przypadkach kombinezon magnetycznego napędu osobistego posiadał będzie także pelerynę przyszytą do tyłu jego rękawów oraz do strony grzbietowej wzdłuż kręgosłupa użytkownika. Peleryna taka pokazana została w części "a" rysunku E4. W przypadku lotów w powietrzu, po rozpostarciu jak skrzydła peleryna ta dostarczy użytkownikowi dodatkowych atrybutów aerodynamicznych podobnych to tych wytwarzanych przez dzisiejsze lotnie. Atrybuty te zwiększą płynność manewrowania i precyzję lotu. Jednak peleryna taka nada również użytkownikowi nieco odpychającego wyglądu błoniastego nietoperza.

E2. Działanie magnetycznego napędu osobistego

Zasada działania magnetycznego napędu osobistego zilustrowana została na **rysunku E3**. Lewa sylwetka tego rysunku pokazuje **siły zewnętrznego** oddziaływania napędu osobistego z polem magnetycznym otoczenia, tj. z polem ziemskim, słonecznym lub galaktycznym. Osiem pędników bocznych zamontowanych w pasie napędu osobistego, zorientowanych zostaje w sposób powodujący ich odpychanie od pola magnetycznego otoczenia. W ten sposób pędniki boczne wytwarzają siły nośne "R" jakie unoszą użytkownika. Miniaturowe pędniki główne zamontowane w podeszwach butów zorientowane są przyciągająco względem pola magnetycznego otoczenia. Dzięki temu formują one dwie siły stabilizujące "A" (tj. prawą "AR" i lewą "AL") jakie wyznaczają położenie wymagane przez użytkownika w czasie jego lotu. Odpowiednie wyważanie wartości sił "R" i "A" powoduje wznoszenie (wzlot), nieruchome zawisanie, lub opadanie użytkownika. Dla przykładu wzlot użytkownika w górę nastąpi gdy wartość sił nośnych "R" przekroczy wartość sił stabilizacyjnych "A". Im większa będzie nadwyżka sił podnoszących "R" nad siłami obniżającymi "A" tym wzlot użytkownika będzie szybszy.

Obie grupy sił "R" i "A" są pochodzenia zewnętrznego, ponieważ powstają one gdy pole magnetyczne otoczenia (np. pole Ziemi) odpycha lub przyciąga pole wytwarzane przez pędniki główne i boczne napędu osobistego. Stąd siły produkowane w wyniku tego oddziaływania z otoczeniem mogą być nazywane "siłami zewnętrznymi". Niezależnie od nich, w napędzie osobistym występuje też i drugi rodzaj sił, jakie mogą zostać nazwane "siłami wewnętrznymi". Są one formowane w wyniku wzajemnych oddziaływań pomiędzy poszczególnymi pędnikami napędu osobistego. Owe "**siły wewnętrzne**" pokazane zostały na prawej sylwetce z rysunku E3. Składają się na nie:

B - Siły wzajemnego odpychania się pędnika głównego z podeszwy jednego buta od pędnika głównego z podeszwy drugiego buta. Siły odpychające "B" wytwarzane są ponieważ bieguny magnetyczne (N, S) w obu pędnikach głównych muszą być zorientowane w tym samym kierunku.

F - Siły wzajemnego odpychania się każdego pędnika bocznego z pasa od innych pędników zawartych w tym samym ośmio-segmentowym pasie. Siły "F" powodować będą odśrodkowe rozprężanie się i napinanie pasa.

Q - Siły wzajemnego przyciągania się pomiędzy każdym z pędników głównych w butach i każdym pędnikiem bocznym w pasie. Owe siły "Q" wytwarzane będą ponieważ bieguny magnetyczne w obu pędnikach głównych muszą zostać zorientowane w kierunku przeciwnym niż bieguny magnetyczne wszystkich pędników bocznych.

Warto zauważyć, że istnieje ściśle podobieństwo pomiędzy siłami "zewnętrznymi" i "wewnętrznymi" panującymi w napędzie osobistym, a podobnymi siłami panującymi w konstrukcji magnokraftu - patrz [1F4.3], lub opis w podrozdziale F4.3 oraz ilustracja z rysunku F15.

Obecność zewnętrznego i wewnętrznego układu sił jest korzystna dla użytkownika napędu osobistego. Oba układy łączą bowiem poszczególne elementy tego napędu w jeden współdziałający zespół. Działanie tego zespołu jest tak dobrane, że przeciwstawne sobie siły nawzajem się balansują. Dla przykładu, kiedy siła nośna "R" i stabilizacyjna "A" starają się rozerwać użytkownika, jednocześnie siły "Q" starają się ścisnąć go wzdłuż tego samego kierunku. W ten sposób ciało użytkownika nie jest ani ściskane ani też rozciągane siłami magnetycznymi. Równocześnie napięty układ wzajemnie przeciwstawnych sobie sił otaczających ciało użytkownika formuje rodzaj "magnetycznego szkieletu" jaki otacza i unosi daną osobę w sposób podobny jak współczesny samochód jest chroniony i unoszony przez jego ramę. Istnieje jednak warunek nałożony na wzajemne balansowanie się sił "R/A" i "Q". Warunek ten stwierdza, że użytkownik nie może zginać nóg w kolanach, ponieważ każde ich zgięcie dostarczyłoby natychmiastowej przewagi siłom "Q" nad siłami "R/A". (Wartość sił Q wzrasta wykładniczo ze zmniejszaniem się odległości pomiędzy pędnikami z butów i pasa.) Jeśli więc użytkownik złamie ten warunek, raz zgięte nogi będą musiały pozostawać

unieruchomione w pozycji kucającej przez resztę lotu. Stąd latanie w pozycji kucania z nogami skrzyżowanymi będzie jedną z dwóch charakterystycznych postaw użytkowników napędu osobistego (druga postawa to nogi sztywno wyprostowane i w rozkroku - jej opis nastąpi). Ciekawe, że latanie z takim właśnie skrzyżowaniem podwiniętych nóg jakby dana osoba siedziała na ubikacji jest często ilustrowane jako typowe zachowanie się średniowiecznych "diabłów" (przykładowo patrz średniowieczna rzeźba diabła znajdująca się na zamku wysokim w Malborku przy wlocie korytarza wiodącego do ubikacji krzyżackich).

Innym warunkiem nałożonym na napęd osobisty to balansowanie sił rozpierających "B" panujących pomiędzy obu pędnikami głównymi zamontowanymi w butach. Siły te utrzymują nogi użytkownika w rozkroku przez cały czas użytkowania napędu osobistego. Wypełnienie więc tego drugiego warunku dostarcza jeszcze jednej postawy charakterystycznej jaka będzie umożliwiała identyfikowanie i odróżnianie użytkowników magnetycznego napędu osobistego (pierwszej generacji). Ich nogi muszą być bowiem trzymane wyprostowane w rozkroku nie tylko podczas lotów przy unieruchomionym ciele, ale także podczas wspomagania tradycyjnych metod poruszania się, takich jak chodzenie, skakanie, pływanie, itp. Przykładowo podczas chodzenia użytkownik tego napędu nie będzie mógł zgiąć w kolanach swoich sztywno wyprostowanych nóg, bowiem każde ich zgięcie w kolanach spowoduje natychmiastowe podwinięcie tych nóg w pozycję kucającą omówioną paragraf wyżej. Stąd często zamiast chodzić jak to się czyni normalnie, użytkownicy tego napędu poruszali się będą długimi skokami odbywającymi się oboma wyprostowanymi i rozkroczonymi nogami naraz bez zginania ich w kolanach (tj. skokami podobnymi do tych wykonywanych przez dzieci imitujące poruszanie się królika). Szokującą dziwność ich postawy i sposobu poruszania się niekiedy pogłębi jeszcze wyjaśniony przy końcu podrozdziału E1 zwyczaj trzymania przez nich swoich rąk wyprostowanych przed sobą jak to czynią lunatycy lub wystawionych w boki od ciała jak policjant regulujący ruch drogowy - patrz też opisy z podrozdziału R4.1. Warto zauważyć, że chociaż poruszanie się z takim ciągłym rozkrokiem nóg i wyciągniętymi do przodu rękami może wyglądać bardzo dziwnie i niezgrabnie, przy bliższym kontakcie zapewne zaimponuje on przypadkowemu obserwatorowi niezwykłą zręcznością, efektywnością i szybkością.

Niezależnie od statycznych oddziaływań siłowych, napęd osobisty wytwarzał też będzie oddziaływania dynamiczne. Oddziaływania te wywoływane będą wirem magnetycznym formowanym wokół pasa użytkownika poprzez 90 przesunięcia fazowe w pulsowaniach pola wytwarzanego przez kolejne pędniki boczne. Wir ten podobny będzie do rotującego pola magnetycznego formowanego przez współczesne silniki asynchroniczne. Dostarczy on napędowi osobistemu szeregu korzystnych atrybutów, przykładowo uformuje on wokół użytkownika rodzaj "pancerza indukcyjnego" jaki spowoduje eksplozyjne odparowanie pocisków wystrzelonych w jego kierunku.

Kolejną grupę oddziaływań w magnetycznym napędzie osobistym pierwszej generacji stanowią oddziaływania konfiguracyjne. Wynikają one z faktu otaczania użytkownika polem magnetycznym o odpowiednim kształcie przestrzennym. Przy właściwym doborze tego kształtu napęd osobisty wytworzy soczewkę magnetyczną jaka odchyli promieniowanie świetlne czyniąc użytkownika niewidzialnym (patrz rysunek F32 i opis z podrozdziału F10.3). Warto tutaj też zauważyć, że użytkownicy napędów osobistych drugiej i trzeciej generacji opisanych w podrozdziałach L4 i M4 niezależnie od stawiania się niewidzialnymi poprzez włączenie owej soczewki magnetycznej, do swej dyspozycji będą także mieli i kilka dalszych sposobów osiągnięcia niewidzialności, takich jak np. wprowadzenie się w stan szybkiego migotania telekinetycznego (patrz opis z podrozdziału L1), manipulacje na upływie czasu (patrz podrozdział M1), czy oddziaływanie na zmysł wzroku obserwatora (patrz modyfikatory wyglądu opisane w podrozdziale N3.2).

Dla sterowania działaniem napędu osobistego używane będą specjalne mikro-komputery. Komputery te odczytywać będą bioprądy odbierane przez specjalnie skonstruowane czujniki umieszczone u podstawy czaszki z tyłu szyi użytkownika. Następnie bio-prądy te tłumaczone będą na sygnały wykonawcze. (Zasada takiego odczytywania i

tłumaczenia omówiona już była w podrozdziałach N1 i N3 - patrz opisy tzw. "interface rozpoznającego myśli" albo "TRI" - podrozdział N3.1.) Stąd wystarczy jedynie pomyśleć o wzlocie do góry, opadaniu czy określonym manewrze, aby zamierzony przez użytkownika rodzaj przemieszczenia został natychmiast zrealizowany bez konieczności wykonywania jakichkolwiek ruchów przez odpowiednie części ciała. Tak samo się stanie w przypadku wspomagania tradycyjnych metod poruszania się, takich jak kroczenie czy skakanie. Wystarczy wtedy jedynie spróbować wykonać dany ruch kończynami aby napęd osobisty wsparł i zwielokrotnił jego moc, zasięg i szybkość. Zasady sterowania napędem osobistym będą przy tym podobne do tych wykorzystywanych w magnokracie. Także metoda wytwarzania wiru magnetycznego będzie taka sama. Jedynie częstotliwość rotowania wiru magnetycznego będzie znacznie wyższa, aby uczynić niemożliwym wytwarzanie wiru plazmowego (który mógłby upalić ręce użytkownika). Jednak nawet gdy prędkość kątowna rotacji wiru magnetycznego będzie zbyt wielka dla zabrania i przyspieszania cząsteczek zjonizowanego powietrza, dysze wylotowe pędników mogą zjonizować powietrze lokalnie. Stąd w nocy powinno być zauważalne wydzielanie się charakterystycznego świecenia wokół pasa i butów tego napędu. Także obce substancje jakie normalnie mogłyby przylegać do kombinezonu (np. błoto czy kurz) w napędzie osobistym będą elektryzowane i odrzucane przez odśrodkową akcję wiru magnetycznego.

E3. Kombinezon z pędnikami głównymi w naramiennikach

W podstawowej wersji napędu osobistego opisanej wcześniej i pokazanej na rysunku E2 pędniki główne wbudowane są w podeszwy butów. Rozwiązanie to wykazuje jednakże poważną wadę jaką jest zbiór sił "B" (patrz prawa część na rysunku E3). Siły te, działając pomiędzy nogami użytkownika, powodują rozpieranie tych nóg w pozycję trwałego rozkroku. W ich więc wyniku ruchy użytkownika muszą odbywać się bez zginania nóg w kolanach i przy ich utrzymywaniu w stałym rozkroku, co w efekcie uniemożliwia osiąganie pełnej swobody i wygody (nie wspominając już że wygląda nieco niezgrabnie, dziwacznie i po prostu śmiesznie).

Aby wyeliminować owe siły "B" zaprojektowana może zostać inna wersja kombinezonu napędu osobistego. Pokazano ją na **rysunku E4** (a). W wersji tej pędniki główne zostały usunięte z podeszw butów i umieszczone w epoletach (1) na ramionach użytkownika. Z punktu widzenia zasady działania napędu takie przemieszczenie pędników z butów do naramienników nie wprowadza większej zmiany. Jednakże dla użytkownika oznacza to uwolnienie nóg od niepożądanych sił i umożliwienie ich swobodnych ruchów. Stąd owa wersja kombinezonu będzie stosowana w każdej sytuacji wymagającej użycia nóg. Jej wadą będzie jednak bliskość do głowy źródeł silnego pola magnetycznego (tj. obu pędników głównych). Stąd twarz i głowa użytkownika w tej wersji muszą zostać osłonięte szczególnie dobrze. W przypadku więc gdy zachodzi niebezpieczeństwo, że warstewka kremu ochraniającego skórę twarzy może zostać starta (np. w wyniku wykonywanej pracy), koniecznym się staje ubieranie przez użytkownika specjalnej maski, opisanej już przy końcu podrozdziału E1.

Oczywiście istnieje możliwość że po nałożeniu na skórę głowy warstewki kremu ochronnego oraz po przesterowanie pola napędu na wydatek nieszkodliwy dla zdrowia (np. na pole niepulsujące takie jak wytwarzane przez magnesy stałe), użytkownicy wersji napędu osobistego z pędnikami w epoletach latali też będą zupełnie bez żadnego nakrycia głowy - np. patrz UFOauta zwany Ausso z rysunku R4. W takim jednak przypadku ich wszystkie włosy na głowie ulegną silnemu naelektryzowaniu i stać będą dęba, sprawiając niesamowite wrażenie na postronnych widzach - patrz istoty nadprzyrodzone opisane w podrozdziale R4.1.

Z powodu tendencji tego napędu do poszerzania ramion użytkownika przez nawzajem odpychające się pędniki główne umieszczone w epoletach (patrz siły "B" w prawej części rysunku E3), ta wersja napędu osobistego nada użytkownikowi charakterystycznego trójkątnego (barczystego) kształtu.

Niekiedy epolety zawierające pędniki główne mogą zostać połączone z pasem za pośrednictwem specjalnych szelek (patrz rysunek R4). Takie skrzyżowane szelki powodują, iż wzajemne położenie pędników głównych i bocznych jest bardziej stabilne i niezawodne. Eliminują one także tendencję do poszerzania ramion użytkownika wspomnianą poprzednio.

W przypadkach gdy wymagane jest wykonanie ciężkich prac fizycznych, omawiana tu wersja napędu może również zostać wzmocniona poprzez dołożenie dodatkowych bransolet z pędnikami wspomagającymi. Na rysunku E4 (a) bransoletki takie oznaczono przez (3).

E4. Wersja napędu osobistego z poduszkami wokół bioder

Kombinezony napędu osobistego mogą zostać poddane dalszym różnorodnym modyfikacjom, dostarczającym im wymaganych zalet operacyjnych. Kolejna z takich zmodyfikowanych wersji pokazana została na rysunku E4 (b). W wersji tej, dłonie użytkownika osłonięte są przed działaniem silnego pola magnetycznego za pomocą specjalnego ekranu. Stąd wersja ta pozwala na wyeliminowanie rękawic używanych w kombinezonie standardowym z rysunku E2. To z kolei umożliwia precyzyjniejsze działania ręczne (np. w celu zmontowania jakiegoś skomplikowanego urządzenia) bez konieczności wyłączania działania pędników napędu.

W opisywanej tu modyfikacji kombinezonu, odpowiednie poduszki osłaniające (1) zostają doszyte wokół ośmio-segmentowego pasa (3). Poduszki te wypełnione są helem, tj. gazem jaki wykazuje najwyższą oporność na zjonizowanie (tj. jaki posiada najwyższy potencjał elektryczny jonizacji). Powierzchnia wewnętrzna obwodu tych poduszek posiada wbudowaną osłonę magnetorefleksyjną (2). Z uwagi na tą osłonę, pole magnetyczne wytwarzane w pędnikach z pasa nie może działać tak silnie na ręce użytkownika jak by ono działało w przypadku użycia kombinezonu standardowego. Stąd w wersji tej ręce nie muszą być osłanianie za pomocą rękawic. Poduszki (1) podzielone są przegrodami (4) na osiem oddzielnych komór podobnych do komór z kołnierza bocznego magnokraftu. Każdy więc pędnik z pasa (3) utrzymywany jest w oddzielnej komorze. To z kolei czyni niemożliwym wytwarzanie wiru plazmowego jaki podążałby za wirem magnetycznym wytwarzanym przez pędniki w pasie. Stąd wyeliminowane jest niebezpieczeństwo upalenia rąk użytkownika. Jednakże kombinezon z poduszkami wokół bioder wygląda bardzo niezwykle, jako że czyni on użytkownika podobnym do gruszki.

Część (b) rysunku E4 pokazuje także alternatywną osłonę głowy użytkownika napędu osobistego. Jest nim przeźroczysty i magnetycznie nieprzenikalny hełm (5) jaki w niektórych modyfikacjach może zastępować kaptur i krem grafitowy z kombinezonu standardowego. Powinno tu być podkreślone, że podobnie jak każde inne nakrycie głowy zilustrowane w niniejszym rozdziale hełm taki może zostać użyty praktycznie z każdą wersją napędu osobistego, nie tylko z wersją z poduszkami wokół bioder.

E5. Osiągi napędu osobistego

Najważniejszym osiągiem napędu osobistego jest to, iż pozwoli on użytkownikowi na latanie w powietrzu bez konieczności odwoływania się do użycia jakiegoś widocznego pojazdu. W ten sposób, będąc jedynie wyposażony w odpowiedni kombinezon który i tak w większości sytuacji musiałby on przywdziawać w celach ubiorczych, użytkownik w każdej chwili może wznieść się w powietrze i przelecieć na dowolny dystans z szybkością ograniczaną jedynie koniecznością oddychania.

Napęd osobisty umożliwia też wspomaganie tradycyjnych sposobów poruszania się (np. skakania, chodzenia, pływania), zwiększając ich zasięg i szybkość, czyniąc je efektywniejszymi niż normalnie, oraz powodując ich funkcjonalne poszerzenie. Dla przykładu pozwala on na: wskakiwanie z poziomu ulicy bezpośrednio na dachy najwyższych budynków,

doganianie w biegu najszybszych samochodów, chodzenie po powierzchni wody, chodzenie po suficie z głową skierowaną w dół, czy wchodzenie na pionowe ściany z poziomym położeniem ciała (jak to czynią owady).

Napęd osobisty posiada też możliwość potęgowania i zwielokrotniania siły fizycznej użytkownika. Możliwość taka rodzi się z użycia pędników wspomagających zamocowywanych jak zegarki na przegubach obu rąk. Wyposażony w nie użytkownik będzie w stanie wywracać całe budynki, wyrwać drzewa z korzeniami, podnosić ogromne głazy, oraz dokonywać wszelkich innych prac wymagających ponadludzkiej siły.

Napęd osobisty dostarczy też swojemu użytkownikowi wielu niezwykle zdolności nieosiągalnych innym sposobem. Jednym z przykładów takich zdolności jest odporność na kule i pociski uzyskiwana dzięki otoczeniu użytkownika wirującym polem magnetycznym formującym rodzaj "pancerza indukcyjnego". Innym niezwykle atrybutem napędu osobistego jest możliwość uczynienia użytkownika niewidzialnym. Niewidzialność taka uzyskiwana jest poprzez uformowanie z pola magnetycznego odpowiedniej soczewki magnetycznej. Soczewka taka odchyła drogę promieniowania świetlnego, czyniąc obecność użytkownika tego napędu niezauważalną dla postronnego obserwatora.

E6. Podsumowanie atrybutów magnetycznego napędu osobistego

Dla cywilizacji tak nisko zaawansowanej jak nasza, która jeszcze nie dorobiła się ani magnokraftu, ani też jego personalizowanej wersji w postaci napędu osobistego, istotna jest znajomość atrybutów magnetycznego napędu osobistego. Wszakże siły konfederacji cywilizacji kosmicznych okupujących naszą planetę, często będą wyposażone w taki właśnie napęd. Z tego też powodu jest istotnym abyśmy nauczyli się rozpoznawać fakt użycia takiego napędu w naszym pobliżu, a także abyśmy byli w stanie identyfikować i rozumieć pochodzenie i możliwości jego najważniejszych atrybutów.

Magnetyczny napęd osobisty dostarcza wykorzystującym go osobom szeregu unikalnych i niezwykle użytecznych cech. Większość z nich podobna jest w efektach do atrybutów magnokraftu, aczkolwiek istnieją też i cechy jakie wystąpią jedynie w przypadku użycia napędu osobistego. Najważniejsze atrybuty magnetycznego napędu osobistego, jakie umożliwiają zidentyfikowanie faktu jego użycia, są jak następuje:

#1. Konieczność ubierania przez użytkowników specjalnego kombinezonu. Kombinezon ten niewiele się różni od ubrań noszonych współcześnie, aczkolwiek obejmuje sobą również nakrycie głowy stąd podstawowa jego wersja przypomina nieco uniformy noszone przez mnichów i zakonnice. Na kombinezon ten składają się:

a) Buty, których podeszwy zawierają wbudowane miniaturowe pędniki główne. Niekiedy pędniki te mogą zostać przeniesione z butów do naramienników.

b) Ośmio-segmentowy pas zawierający pędniki boczne.

c) (Niekiedy) Dwie bransoletki zakładane na nadgarstki (lub nawet zamocowywane do zewnętrznych powierzchni rękawic przywdziewanych na dłonie) a zawierające pędniki wspomagające jakie dopomagają w ciężkich pracach fizycznych. Owe pędniki wspomagające nie są wykorzystywane podczas lotów, stąd są one zakładane tylko jeśli wymagane jest zwiększenie siły fizycznej użytkownika.

d) Komputer sterujący zamontowywany w tylnej części szyi.

e) Jednocześnieowy kombinezon obejmujący także kaptur lub hełm.

f) (Niekiedy) Pelerynę przyszytą do rękawów i strony grzbietowej kombinezonu. W przypadku lotów w powietrzu, po rozpostarciu jak skrzydła peleryna ta dostarczy użytkownikowi dodatkowych atrybutów aerodynamicznych (podobnych to tych wytwarzanych przez dzisiejsze lotnie) jakie zwiększą płynność jego manewrów. Równocześnie jednak nada mu nieco odpychającego wyglądu błoniastego nietoperza.

g) Rękawice z błono-podobnymi łącznikami pomiędzy palcami.

h) Krem na bazie grafitu do pokrywania fragmentów skóry użytkownika wystawionej na działanie silnego pola magnetycznego (np. jego twarzy). W niektórych wypadkach krem ten może zostać zastąpiony maską ochronną na twarz podobną do masek używanych przez rabusiów bankowych.

#2. Konieczność przyjmowania przez ciało szczególnej postawy charakteryzowanej przez: sztywno wyprostowane nogi utrzymywane w ciągłym rozkroku (lub zakrzywione w pozycję kucającą), dłonie odepchnięte daleko od pasa, palce rozcapierzone, itp. Ponadto, wszystkie włosy "stojące dęba" z powodu działania ładunków elektrycznych gromadzących się na powierzchni użytkownika. W niektórych przypadkach użytkownik tego napędu poruszał się też będzie z rękami wyciągniętymi do przodu lub rozpostartymi z boków ciała. Powyższe czyni postawę i ruchy użytkownika napędu osobistego wyglądające bardzo nienaturalnie i niezręcznie (aczkolwiek w akcji mogące zadziwić swą efektywnością i zręcznością).

#3. Powodowanie jarzenia się oczu a niekiedy również i zębów użytkownika. Silne fluorescencyjne świecenie się tych oczu i zębów jest powodowane przez oddziaływanie na nie potężnego pola magnetycznego (podobnie jak niektóre odmiany niewidzialnego promieniowania elektromagnetycznego powodują jarzenie się oczu u oświetlonych nim zwierząt). Z kolei takie silne jarzenie się oczu i/lub zębów, często akompaniowane przez świecenie się kombinezonu napędu osobistego, nadaje użytkownikowi tego napędu bardzo niezwykłego, "nadprzyrodzonego" wyglądu.

#4. Zdolność osoby ubierającej taki napęd do dokonywania działań zaprzeczających prawom fizyki. Przykładowo jest ona zdolna do bezgłośnego latania w powietrzu połączonego z możliwością przyjmowania dowolnej orientacji ciała niezależnej od panujących sił przyciągania grawitacyjnego (np. oprócz zwykłego stania na podłodze, także możliwości zwisania z sufitu, poziomego stania na ścianach, pochylania się pod dowolnym kątem, itp.). Sterowanie pozycją zajmowaną przez ciało nie wymaga dokonywania żadnego ruchu, bowiem odbywa się poprzez przetwarzanie przez komputer kontrolujący bioprądy pobieranych bezpośrednio z szyi użytkownika.

#5. Zdolność do wspomagania normalnych sposobów poruszania się (np. skakania, chodzenia, pływania). Takie wspomaganie czyni możliwym dokonywanie poruszeń jakie zdają się zaprzeczać naszemu obecnemu rozumieniu praw natury i możliwości człowieka, dla przykładu:

- a) Spacerowania po suficie w pozycji do góry nogami.
- b) Kroczenia w górę lub dół pionowych ścian z ciałem przymującym pozycję poziomą (tj. w sposób jak to czynią owady).
- c) Dokonywania skoków na ogromne wysokości i odległości (np. wskakiwanie z chodników bezpośrednio na dachy okolicznych budynków). Skoki te odbywały się będą albo przy sztywno wyprostowanych nogach i to bez jakiegokolwiek ich zginania w kolanach, albo nawet kiedy osoba je dokonująca wygląda jakby kucała.
- d) Spacerowanie po powierzchni wody.

#6. Niezwykłe zdolności nabywane przez nosicieli napędu osobistego. Wymieńmy niektóre z nich:

- a) Odporność na kule wynikająca z zabezpieczającego działania "pancerza indukcyjnego".
- b) Czynienie się niewidzialnym poprzez włączenie działania "soczewki magnetycznej" jaka ugina drogę światła.
- c) Poruszanie się (latanie, chodzenie, lub skakanie) z prędkościami ograniczonymi jedynie przez funkcje fizjologiczne ciała użytkowników (głównie przez możliwość oddychania). Takie szybkie poruszanie się nie wymaga przy tym użycia żadnego widocznego wehikułu.

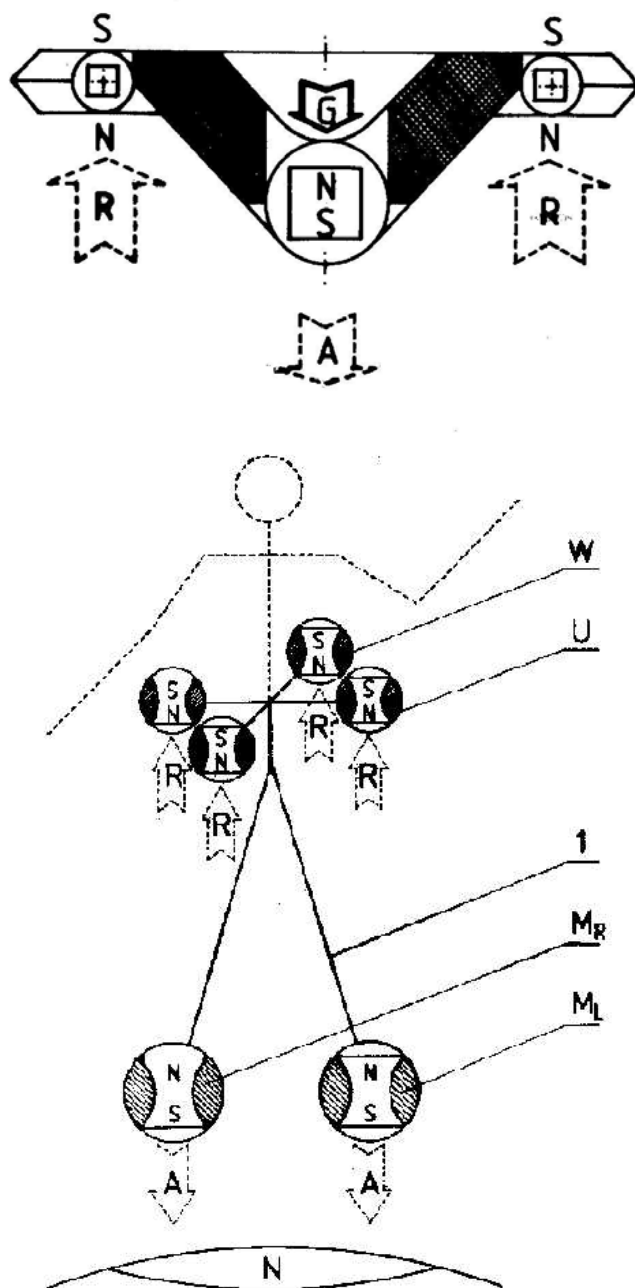
d) Niezwykła siła fizyczna uzyskiwana dzięki akcji pędników wspomagających. Siła ta pozwala na wywracanie budynków, wrywanie drzew z korzeniami, podnoszenie ogromnych skał, oraz wykonywanie innych prac jakie na obecnym poziomie naszej wiedzy wyglądają jako wymagające sił nadprzyrodzonych.

e) Biologiczna sterylizacja otoczenia poprzez zabijanie mikroorganizmów jakie znajdują się w zasięgu potężnego pola napędu (sterylizacja ta z kolei może spowodować różnorodne konsekwencje biologiczne).

#7. Wymuszanie magnetycznych zmian w otoczeniu, podobnych do zmian powodowanych przez napęd magnokraftu. Ich przykłady będą obejmować: magnetyczne wypalanie śladów pod podeszwami butów zaopatrzonych w pędniki, odrzucanie od kombinezonu substancji i obiektów z otoczenia (np. błota i kurzu), elektryzowanie materiałów izolacyjnych (np. włosów które u użytkowników tego napędu zawsze będą stały dęba), jonizowanie powietrza (jakie przy przyciemnionym świetle może się nawet jarzyć w okolicach pasa i butów), produkcja aktywnego ozonu którego zapach będzie towarzyszył użytkownikowi tego napędu (zapach ten przez osoby nieorientowane może zostać pomyłony z zapachem siarki), itp.

* * *

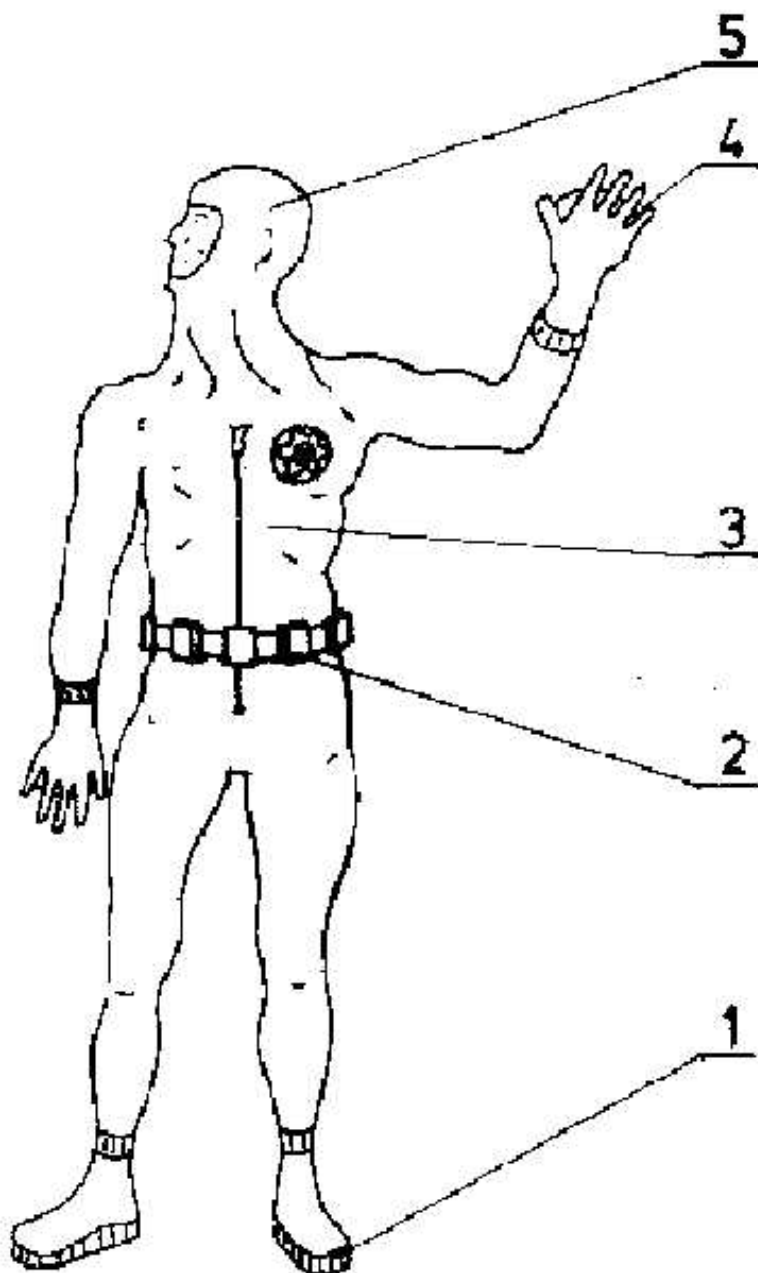
W uzupełnieniu do powyższych informacji warto tu uwypuklić, iż użytkownicy magnetycznego napędu osobistego pierwszej generacji uzyskują wszystkie swoje nadzwyczajne zdolności **tylko w przypadku gdy ubierają oni kombinezon** tego napędu oraz gdy jego pędniki pozostają włączone. Praktyczna (i jedyna) więc metoda obezwładniania użytkowników tego napędu polegałaby na radzeniu sobie z nimi tylko w chwili gdy z jakiegoś tam powodu nieopatrznie zdjęli oni ten kombinezon (np. dla zażycia kąpieli). Oczywiście po jego zdjęciu użytkownicy ci stają się tak samo podatni na wszelkie niebezpieczeństwa jak wszyscy inni śmiertelnicy, zaś ich jedyną obronę zaczyna stanowić zwykła siła, sprawność i trening fizyczny oraz posiadana inteligencja (patrz też podrozdział R4).



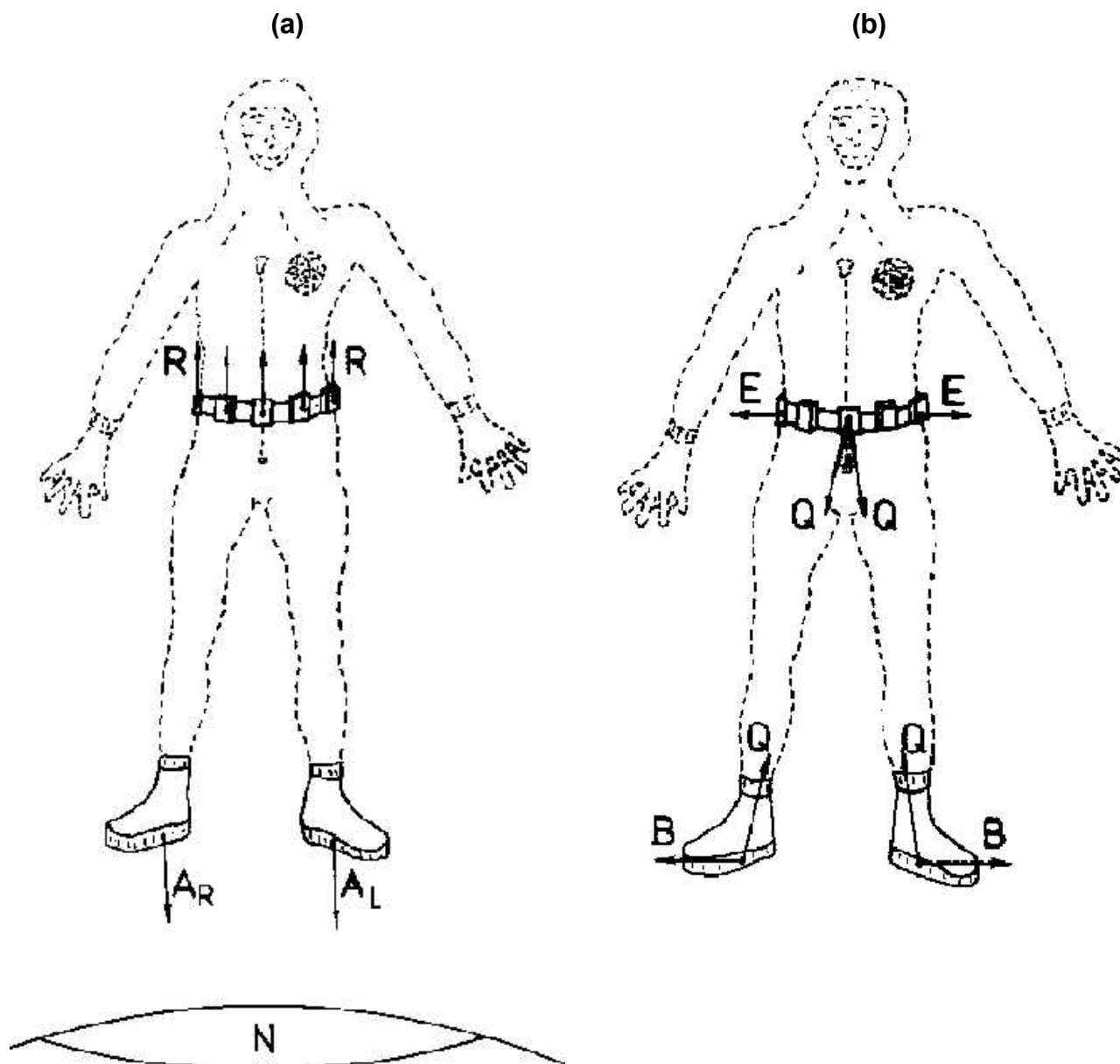
Rys. E1. Podobieństwa pomiędzy układami napędowymi magnokraftu oraz magnetycznego napędu osobistego.

(**Góra**) Magnokraft typu K3 lecący w pozycji "wiszącej" (porównaj tą pozycję z pozycją "stojącą" zilustrowaną na rysunku F4).

(**Dół**) Układ napędowy magnetycznego napędu osobistego, formujący strukturę odpowiadającą kształtowi sylwetki ludzkiej. Układ ten dostarcza zasady dla wyprodukowania kombinezonu magnetycznego napędu osobistego. Osoby wyposażone w taki napęd będą mogły latać w powietrzu bez użycia jakiegokolwiek widocznego wehikułu. Pokazany tu układ składa się z ośmiu pędników bocznych (oznaczonych U, V, W, X) zamontowanych do wnętrza odpowiednio skonstruowanego pasa. Pędniki te wytwarzają siłę nośną (R) (od angielskiego "repulsion"). Ponadto, układ posiada też dwa miniaturowe pędniki główne (oznaczone jako M_R , M_L) zamontowane w podeszwach prawego i lewego buta. Wytwarzają one siły stabilizujące (A) (od angielskiego "attraction forces"). Ciało (1) użytkownika tego napędu dostarcza "szkieletu nośnego" (albo "ramy") jaka wiąże te pędniki w jeden kooperujący system ale jednocześnie utrzymuje je wszystkie w wymaganej odległości od siebie.



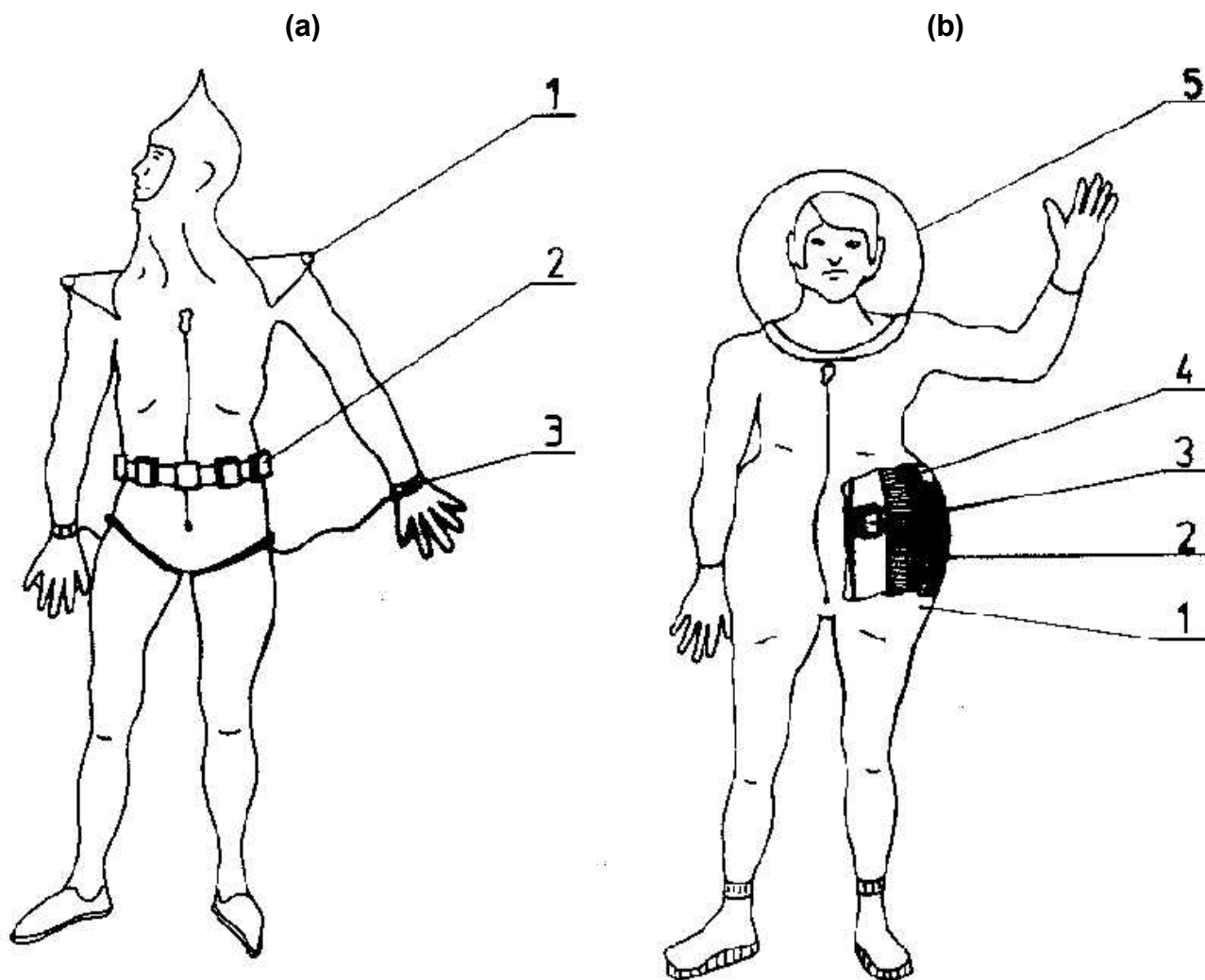
Rys. E2. Standardowy kombinezon napędu osobistego. Pokazano jego ogólny wygląd i elementy składowe. Użytkownicy takiego napędu będą w stanie bezgłośnie latać w powietrzu, spacerować po powierzchni wody, wykazywać odporność na działanie broni palnej, stawać się niewidzialnymi, itp. Na kombinezon ten składają się: (1) buty których podeszwy zawierają wmontowane pędniki główne; (2) ośmio-segmentowy pas zawierający pędniki boczne; (3) jednoczęściowy kombinezon wykonany z materiału magnetorefleksyjnego, jaki obejmuje także kaptur (5) lub hełm; (4) rękawice z błonopodobnymi łącznikami pomiędzypalcowymi. Wszystko to uzupełnione jest kremem na bazie grafitu jaki okrywa odsłonięte części skóry dla zabezpieczenia ich przed działaniem silnego pola magnetycznego, oraz komputerem kontrolnym zamocowywanym z tyłu szyi, jaki odczytuje bioprądy użytkownika i zamienia je na działania napędowe. Kiedy cięższa praca musi zostać wykonana, dodatkowe bransoletki zawierające pędniki wspomagające mogą być nakładane na przeguby rąk (pokazane jako (3) na rysunku E4 "a"). Pędniki te kooperują z pędnikami z pasa i butów, dostarczając użytkownikowi napędu "nadprzyrodzonej" siły fizycznej, np. umożliwiającej mu wrywanie dębów z korzeniami, unoszenie ogromnych głazów, powalanie budynków, itp.



Rys. E3. Siły magnetyczne formowane przez napęd osobisty. Pokazano układ sił zewnętrznych (patrz lewa sylwetka) i wewnętrznych (patrz prawa sylwetka). Należy zauważyć że oba te układy sił neutralizują się nawzajem. Podczas gdy siły "R" i "A" działające w przeciwnych kierunkach rozciągają ciało użytkownika, siły "Q" równocześnie je ściskają. Jedynie siły "B" pozostają niezrównoważone, utrzymując nogi użytkownika w ciągłym rozkroku ułatwiającym identyfikację faktu użycia tego napędu.

(lewa - „a”) Układ **sił zewnętrznych** formowany wskutek oddziaływania pędników napędu osobistego z polem magnetycznym otoczenia. Siły te obejmują: R - siłę nośną wytwarzaną w wyniku oddziaływań odpychających z polem otoczenia; A - siły stabilizacyjne wytwarzane w efekcie oddziaływań przyciągających. Indeksy: R - (right) prawa, L - (left) lewa.

(prawa - „b”) Układ **sił wewnętrznych** formowanych w wyniku oddziaływania pędników magnetycznych pomiędzy sobą. Siły te obejmują: B - siły wzajemnego odpychania się od siebie obu pędników głównych (powodują one stałe odseparowanie (rozkrok) nóg użytkownika); F - siły wzajemnego odpychania się od siebie pędników bocznych (powodują one rozprężanie się pasa); Q - siły wzajemnego przyciągania powstające pomiędzy każdym pędnikiem głównym i każdym pędnikiem bocznym (jeśli wytracone z równowagi poprzez zgięcie nóg, siły te powodują latanie użytkownika w pozycji przykucniętej ze skrzyżowanymi nogami).



Rys. E4. Modyfikacje standardowego napędu osobistego. Pokazane tu przykłady dwóch takich modyfikacji wcale w rzeczywistych napędach nie muszą zostać użyte w dokładnie takim samym zestawie.

(a) Wersja napędu osobistego z pędnikami głównymi zamontowanymi w epoletach. Pokazane zostały: (1) jeden z dwóch pędników głównych; (2) ośmio-segmentowy pas zawierający pędniki boczne; (3) jedna z dwóch bransoletek wspomagających zakładanych na przeguby rąk (niekiedy mogą one też przyjmować formę kwadratowych płytek naszywanych na górnej powierzchni rękawic użytkownika). Bransoletki te zawierają dodatkowe pędniki wspomagające (nie używane do lotów) jakie zwielokrotniają siłę fizyczną użytkownika kiedy musi on dokonywać jakiejś pracy fizycznej. Rysunek ukazuje też skrzydłopodobną pelerynę przyszytą do skafandra wzdłuż kręgosłupa i tyłu rękawów, jaka aerodynamicznie powiększa płynność lotów (jak współczesna lotnia). Dla zwiększenia wytrzymałości i stateczności tego kombinezonu, czasami dwie skrzyżowane szelki wzmacniające będą dodatkowo łączyły pas z epoletami (patrz rysunek R4).

(b) Wersja napędu osobistego z przezroczystym hełmem oraz ochronną poduszką wokół bioder. Pokazane zostały: (1) poduszki otaczające biodra jakie chronią ręce użytkownika napędu przed działaniem silnego pola magnetycznego i elektrycznego; (2) magnetycznie nieprzenikalny ekran oraz przeciw-elektrostatyczna izolacja montowane na zewnętrznych powierzchniach poduszek; (3) jeden z segmentów ośmio-segmentowego pasa zawierającego pędniki boczne; (4) jedna z przedziałek jaka dzieli poduszkę na osiem oddzielnych komór (każda z tych komór izoluje jeden pędnik boczny); (5) magnetycznie nieprzenikalny hełm osłaniający głowę.