

DZIAŁANIE PRĄDÓW ELEKTRYCZNYCH NA USTRÓJ

Dodatkowo jak i ujemne skutki, jakie powoduje prąd elektryczny przepływający przez organizm ludzki, zawsze były w kręgu zainteresowania człowieka. Prace nad wpływem prądu na organizm zaczęły się zaraz po zbudowaniu przez Alessandro Volta pierwszego ogniwa galwanicznego w 1800 roku. Sam Volta badał wpływ bodźców elektrycznych na zmysł wzroku i smaku. (Najbardziej chyba powszechnie znanym, ujemnym skutkiem działania prądu jest porażenie elektryczne). Dopiero jednak w ostatnich kilkudziesięciu latach wiedza na ten temat poszerzyła się znacznie. Badania dokonywane w wielu krajach, także w Polsce, odpowiednio ją usystematyzowały.

Niniejszy artykuł ma na celu przybliżenie czytelnikowi ogólne działanie prądu na komórkę ludzką i zasygnalizowanie obszerności tematu, nie wgłębiając się w tematykę prądów selektywnych, opisywanych w innych publikacjach dostępnych na rynku.

Rodzaje prądów elektrycznych.

W fizjologii i w medycynie stosuje, się w różnych celach prądy elektryczne różnego rodzaju np.:

Prąd stały (galwaniczny) - jest to prąd jednokierunkowy, przeważnie niskiego napięcia (do 80 woltów)

Prąd stały przerywany (galwaniczny przerywany) - jest to prąd jednokierunkowy niskiego napięcia, przerywany systematycznie z częstotliwością od 30 razy na minutę do 120 razy na minutę (kształt impulsu prostokątny).

Prąd stały przerywany o regulowanym kształcie impulsu (dla prądów selektywnych jest to kształt trapezu) o częstotliwościach rzędu kilku Hz.

Prąd zmienny sinusoidalny (sinus-farad) -jest to prąd zmienny o częstotliwości 50-90 Hz (poszczególne impulsy trwają około 0,02-0,01 s); prąd taki otrzymuje się z sieci oświetleniowej, a w badaniach fizjologicznych otrzymuje się z układu silnik-prądnica, albo też z generatora prądowego.

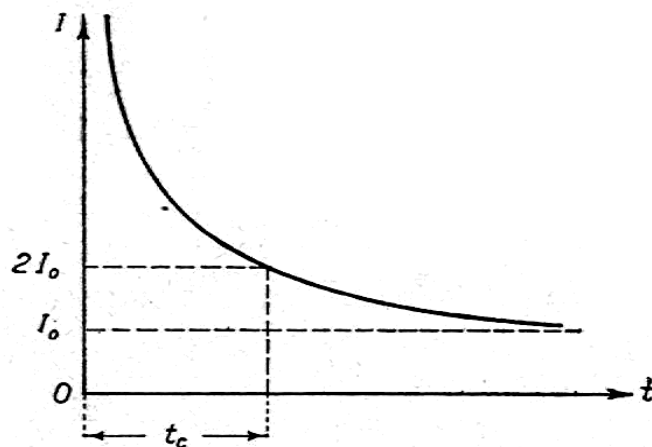
Prąd faradyczny normalny jest to prąd zmienny niskiego napięcia, asymetryczny przerywany o częstotliwości rzędu kilkuset Hz (poszczególne impulsy trwają około 0,001 s); prąd taki otrzymuje się najczęściej z cewki indukcyjnej Ruhmkorfa

Prąd faradyczny modulowany - jest to prąd faradyczny o zmienianej stopniowo amplitudzie, a czasem o zmienianej częstotliwości impulsów.

Prąd diatermiczny - jest to prąd zmienny o wielkiej częstotliwości rzędu 500 kHz do 100 MHz, tzn. o długości fali od 600 metrów do 3 metrów; najczęściej stosuje się prąd diatermiczny długofalowy (długości fali około 300 metrów) i prąd diatermiczny krótkofalowy (długości fali od 3 do 12 metrów)

Najmniejszy czas działania prądu- stałego; potrzebny do spowodowania reakcji fizjologicznej, nazywamy czasem użytecznym. Czas użyteczny służy za miarę pobudliwości: im czas jest krótszy, tym pobudliwość jest większa.

Ponieważ bezpośrednie badania progu pobudliwości dają wyniki trochę rozbieżne, przeto dla określenia pobudliwości nerwu i mięśni zostało wprowadzone w fizjologii i w medycynie pojęcie tzw. chronaksji. Określa się je mianowicie w ten sposób, że wyznacza się najpierw minimalna wartość prądu stałego, dającą pobudzenie (wartość ta nosi nazwę reobazy), a następnie podwaja się ją i dla tej podwojonej wartości prądu wyznacza się jego czas użyteczny. Ten czas nosi nazwę chronaksji. Na ryc.2 pokazany jest wykres zależności wartości natężenia prądu stałego od jego czasu użyteczności; I_0 oznacza reobazę, $2I_0$ - prąd chronaksji, t_c - chronaksję.



Ryc. 2. Zależność natężenia prądu od czasu jego użyteczności. I oznacza reobazę, I_0 oznacza prąd chronakcji, a t_c oznacza chronakcję.

Prądy wielkiej częstotliwości wykazują nawet przy dużym napięciu i dużej gęstości słabe działanie elektrochemiczne, a to z powodu częstej zmiany kierunku napięcia; jony elektrolityczne mają dużą masę, a więc i dużą bezwładność; zyskują wobec tego małe przyspieszenie i nie mogą podczas trwania napięcia jednego kierunku przesunąć się dostatecznie daleko, żeby spowodować wyraźną zmianę koncentracji. Dlatego też; końce wtórnego uzwojenia transformatora Tesli można trzymać w rękach bez szkody dla organizmu, chociaż amperomierz, włączony szeregowo w obwód z ciałem, wskazuje przepływ prądu rzędu kilku miliamperów i chociaż żarówki włączone do tego obwodu świecą jasno.

Dawniej przypuszczano, że efekt ten pochodzi stąd, że prądy wielkiej częstotliwości przepływają w tym przypadku tylko po powierzchni człowieka (tzw. efekt naskórkowy). Pogląd ten jednak nie jest słuszny, gdyż efekt tego rodzaju mogą dawać tylko bardzo dobre przewodniki, jak metale, ale nie takie, jakimi są organizmy.

Mechanizm działania prądów elektrycznych. Napięcia i prądy czynnościowe.

Tkanki ciała ludzkiego składają się z substancji elektrolitycznych o oporze właściwym rzędu kilkuset Ω m. Tkanki zawierają jony różnych soli, przeważnie jony Na (dodatnie) i jony Cl (ujemne) - w ilości około 6 g na 1l, a poza tym zawierają już w znacznie mniejszej ilości jony: K, Ca, Mg, P i inne.

W tkankach rozłożone są równomiernie komórki zamknięte półprzenikliwą błoną, której opór elektryczny jest znacznie większy od oporu otaczającej je substancji. Wewnątrz komórek znajduje się substancja elektrolityczna o podobnych właściwościach elektrycznych, jak substancja zewnętrzna. Kształt komórek można pominąć początkowo przy ogólnych rozważaniach, podobnie jak i strukturę wewnętrzną komórki. Opór elektryczny pojedynczej komórki jest spowodowany prawie wyłącznie przez opór błony komórkowej. Opór elektryczny tkanki dla prądu stałego jest określony głównie przez opór elektryczny pojedynczej błony komórkowej i przez liczbę komórek przypadających na 1 cm^3 tkanki. Wartość oporu elektrycznego tkanki dla prądu stałego nie ma na ogół większego znaczenia praktycznego, stanowi ona natomiast ważną podstawę dla badania mechanizmu przewodzenia prądów zmiennych.

Pod wpływem przyłożonego napięcia zewnętrznego jony przesuwają się w polu elektrycznym wewnątrz komórek i w przestrzeniach międzykomórkowych, wskutek czego w poszczególnych częściach tkanki przepływa prąd elektryczny i powstaje zmiana koncentracji jonów, wywołująca polaryzację elektrolityczną i miejscowe różnice potencjałów elektrycznych. Zjawiska te mają duże znaczenie fizjologiczne, a w szczególności mogą wywołać podrażnienia tkanki, jeżeli tylko podnieta elektryczna przekroczy pewną charakterystyczną dla danej tkanki wartość, tzw. próg pobudliwości.

Dla komórek nerwowych i mięśniowych, które mają kształt walcowy, zjawisko

przewodnictwa elektrycznego na poszczególnych odcinkach przebiega podobnie jak dla kabli elektrycznych. Spadek potencjału przy przepływie prądu wzdłuż komórki następuje wtedy zgodnie z prawem Ohma, a spadek potencjału w poprzek komórki można pominąć przy rozpatrywaniu oporu całej tkanki.

Inaczej przedstawia się sprawa przewodzenia przez komórki i tkanki prądu zmiennego, mianowicie w tym przypadku należy uwzględnić opór pojemnościowy i opór samoindukcyjny komórek, podobnie jak przy obliczaniu zawady elektrycznej. Opory te zależą od częstotliwości kołowej co i tylko w niektórych przypadkach, kiedy pojemność komórki i jej samoindukcja są tak dobrane, że opór zawady znika, zagadnienie bardzo się upraszcza i opór pozorny komórki zbliża się do oporu omowego.

Zmiany koncentracji jonów są na ogół proporcjonalne do gęstości prądu (tzn. do natężenia prądu przypadającego na jednostkę powierzchni ciała), dlatego największe zmiany fizjologiczne zachodzą w tych miejscach, w których elektrody przyłożone do ciała mają powierzchnie najmniejszą (przy tym samym natężeniu prądu).

Koncentracji jonów przeciwdziała zjawisko dyfuzji jonów, dzięki któremu jony dążą do takiego położenia, przy którym rozkład ich jest najbardziej równomierny

Zjawisko dyfuzji zachodzi w tym większym stopniu, im większe są różnice koncentracji jonów i wymaga pewnego czasu dla wyrównania rozkładu jonów; z tego względu prądy krótkotrwałe (przy

których dyfuzja jonów gra rolę znikomą) wykazują znacznie silniejsze działanie drażniące niż prądy długotrwałe, przeprowadzające przez ciało tę samą ilość naboju elektrycznego, tym się tłumaczy np. fakt, że prądy otwarcia w induktorze działają pod względem fizjologicznym dużo silniej niż prądy zamknięcia.

Działanie elektropatologiczne prądów elektrycznych

Najbardziej niebezpieczne dla życia są prądy zmienne o małej częstotliwości, rzędu 30-150 okresów na sekundę. Górna granica bezpieczeństwa wynosi dla nich około 200 woltów napięcia; dla prądów stałych granica ta przesuwą się do 500 woltów. Jednak przy nieszczęśliwym zbiegu okoliczności mogą zająć przypadki śmiertelne nawet przy 60 woltach prądu zmiennego i 250 woltach prądu stałego i dlatego w technice przyjmuje się jako graniczną wartość bezpieczeństwa napięcie 42 woltów prądu zmiennego.

Wypadki porażenia prądem mogą się zdarzyć np. przy dotykaniu przewodów mokrymi rękami, szczególnie w wannie (jeśli ta łączy się przypadkowo z jednym biegunem elektrycznym, przy dotknięciu ręką drugiego bieguna itp.}. Granice bezpieczeństwa są dla różnych osób różne, poza tym nagle, nieoczekiwane uderzenia elektryczne działają dużo silniej niż uderzenia spodziewane.

W większości wypadków śmierć następuje prawdopodobnie przez zatrzymanie czynności serca wskutek drżeń włókienkowych. Przy dłuższym działaniu (ponad 2 minuty) następuje uduszenie wskutek ogólnego skurczu tężcowego mięśni. Często jednak można jeszcze uratować porażonego prądem przez zastosowanie sztucznego oddychania.

Prądy wielkiej częstotliwości nie są niebezpieczne nawet przy napięciu 100000 woltów i więcej. Działanie fizjologiczne jest w przybliżeniu odwrotnie proporcjonalne do pierwiastka kwadratowego z częstotliwości (w pewnym zakresie częstotliwości).

Do ujemnych skutków działania prądu należy też rozmnażanie się komórek rakowych w organizmie, dlatego osoby będące na chemioterapii nie powinny stosować prądów w leczeniu uzupełniającym.

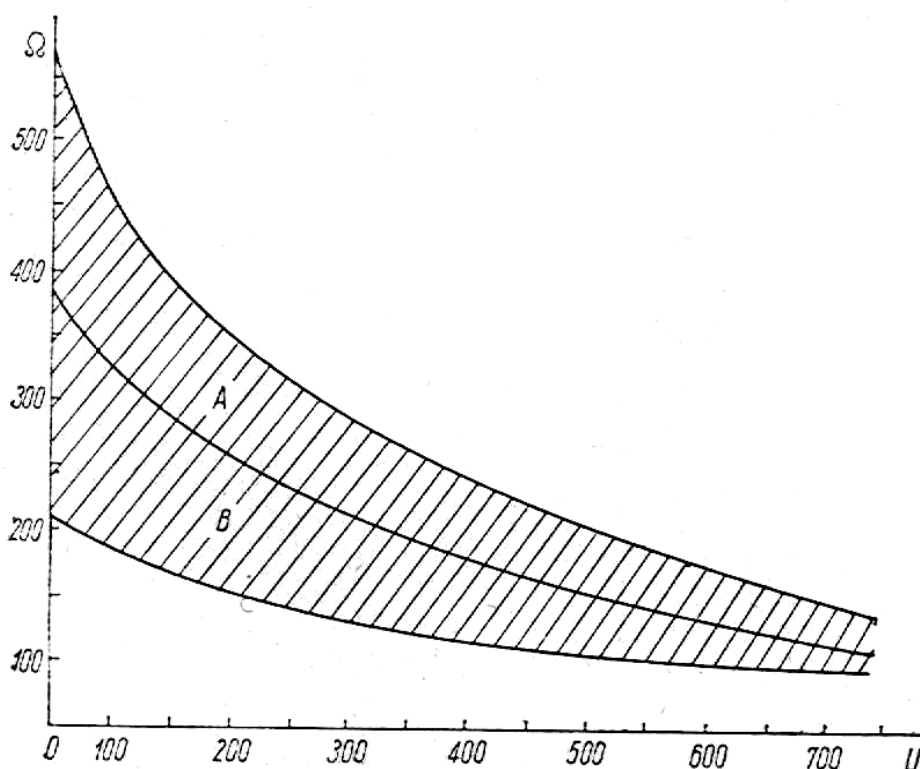
Opór elektryczny ciała ludzkiego dla prądu stałego.

Mechanizm tych zjawisk jest na ogół bardzo różny w zależności od rodzaju nośników ładunku elektrycznego w danej substancji (elektronów swobodnych, jonów) jak również od rodzaju i wielkości przyłożonej siły elektromotorycznej (różnicy potencjałów), od temperatury itd.

W tkankach żywego organizmu pod wpływem przyłożonego napięcia zewnętrznego będą

występowały najczęściej zjawiska typu elektrochemicznego ze względu na to, że, tkanki składają się w dużej mierze z różnego rodzaju elektrolitów, poprzedzielanych nieprzewodzącymi elektrycznie błonami. Zresztą o przewodnictwie elektrycznym tkanek decyduje najczęściej ilość zawartej w nich wody. Niezależnie jednak od tego mogą wchodzić w grę jeszcze inne zjawiska natury elektrycznej, które należy wziąć pod uwagę w poszczególnych przypadkach.

Przy pomiarach prądem stałym otrzymuje się dla całego ciała ludzkiego przy doprowadzeniu różnicy potencjałów do rąk jak ma to miejsce w przypadku stosowania prądów selektywnych, ewentualnie do jednej ręki i do przeciwległej nogi, opór elektryczny wartości rzędu paru tysięcy omów. Na rys.1 pokazana jest zależność oporu elektrycznego dla ciała ludzkiego od przyłożonego napięcia (według Freiberga), przy czym wartości oporu mogą się zmieniać w granicach pomiędzy górną i dolną krzywą w obszarze zakreskowanym w zależności od wilgotności ciała.



Ryc. 1. Zmiana oporu ciała ludzkiego w zależności od napięcia przyłożonego:
A — dla ciała suchego, B — dla ciała mokrego.

Opór właściwy (ρ) tkanki jest funkcją kilku różnych parametrów takich, jak: opór właściwy płynu zewnątrzkomórkowego (ρ_1), opór właściwy zawieszonych w niej komórek (ρ_2), względna objętość komórek (φ), opór właściwy plazmy (ρ_3), opór właściwy błony komórkowej (ρ_4), promień komórki (a), i czynnika geometrycznego (f) określającego kształt komórki.

Z rozważań teoretycznych i doświadczalnych nad różnymi rodzajami zawiesin w roztworach elektrolitycznych ustalono następującą zależność dla ρ :

$$\rho = \rho_1 \frac{(1-\varphi)\rho_1 + (f+\varphi)\rho_2}{(1+f\varphi)\rho_1 + f(1-\varphi)\rho_2} \quad [1]$$

Przy tym $\rho_2 = \rho_3 + \rho_4/a$ czynnik $\varphi = v_K/V$ określa stosunek objętości komurek (v_K) do

całkowitej objętości tkanki (V) a czynnik geometryczny f przyjmuje wartości: 1.5 dla komórek kulistych.

Wartości liczbowe poszczególnych oporów właściwych są rzędu: ($\rho_1 \sim 10^2 \Omega m$, $\rho_2 \sim 10^4 \Omega m$, $\rho_3 \sim 1 \Omega m$, $\rho_4 \sim 0,1 \Omega m$ (przy $a = 10^{-5} m$). Jak łatwo zauważyć ze wzoru [1], podstawiając różne wartości na ρ_2 , ich wpływ na wartość ρ_1 jest stosunkowo nieduży (rzędu paru procent) tak, że można napisać wzór [1] w postaci uproszczonej

$$\rho = \rho_1 \frac{f + \varphi}{f(1 - \varphi)}, \quad [2]$$

W poniższej tabeli podane są wartości liczbowe oporu właściwego różnych tkanek.

Tkanka	Opór właściwy Ohm
Płyn mózgowo- rdzeńiowy	0
Surowica krwi	0
Mięśnie	1, 52
Krew	0
Wątroba	0
Mózg	25
Tkanka tłuszczowa	50
Skóra sucha	3030
Kości bez okostnej	2×10^2

Opór elektryczny ciała ludzkiego dla prądów zmiennych niskiej częstotliwości

W przypadku, kiedy do tkanki przyłożymy zmienne napięcie U o ustalonej częstotliwości kątowej ω , przez tkankę popłynie prąd zmienny o natężeniu I, takim że

$$U = ZI, \quad [3]$$

gdzie Z oznacza opór pozorny tkanki (tzw. impedancję) i równa się

$$Z = \sqrt{R^2 + \left(\omega L - \frac{1}{\omega C} \right)^2} \quad [4]$$

Jeżeli impedancja dotyczy tylko błony komórkowej (Z_m), a wewnątrz komórki jest wolne od impedancji, to można napisać

$$Z_2 = r_2 + \frac{Z_m}{a}, \quad [5]$$

i wtedy

$$Z = r_1 \frac{(1-\rho)r_1 + (2+\rho)\left(r_2 + \frac{Z_m}{a}\right)}{(1+2\rho)r_1 + 2(1-\rho)\left(r_2 + \frac{Z_m}{a}\right)} \quad [6]$$

Jeżeli błona komórkowa ma pojemność C_m na jednostkę powierzchni i nie ma oporu upływności, to przy małych częstotliwościach otrzymamy: tylko opór r_0 :

$$r_0 = r_1 \frac{1-\rho}{1+2\rho}, \quad [7]$$

a przy dużych częstotliwościach

$$r_\infty = r_1 \frac{(1-\rho)r_1 + (2+\rho)r_2}{(1+2\rho)r_1 + 2(1-\rho)r_2} \quad [8]$$

Pomiary r_1 i r_0 określają wielkość ρ , a pomiar r_∞ określa wielkość r_2 , wewnętrzny opór właściwy komórki.

Przy pomiarach osiowego oporu elektrycznego pojedynczego włókna otrzymuje się dla bardzo dużych częstotliwości

$$R_\infty = \frac{r_1 r_2}{r_1 + r_2} s, \quad [9]$$

gdzie s oznacza odległość pomiędzy elektrodami.

Dla $\omega=0$ wyrażenie na opór (R_0) składa się z dwóch składników - jednego podanego we wzorze [9] i drugiego bardziej skomplikowanego.

Doświadczalnie mierzy się R_∞ , R_0 przy różnych odległościach s i stąd wyznacza się r_1 , r_2 i r_4 i w ten sposób określa opory: wewnętrzny, zewnętrzny i błony komórkowej.

Dla częstotliwości pośrednich otrzymuje się na impedancję wzór:

$$Z = R_\infty + (R_0 - R_\infty) \sqrt{Z_m / r_4} \quad [10]$$

Stały prąd elektryczny, płynący pomiędzy elektrodami przyłożonymi do skóry zwierzęcia, napotyka w pierwszym rzędzie na opór elektryczny skóry. Opór reszty tkanek gra już wtedy rolę drugorzędna. W skórze suchej prąd płynie głównie przez kanaliki gruczołów potowych i niektórych tłuszczowych.

W skórze wilgotnej przewodnictwo elektryczne jest znacznie większe. Na skutek przesunięcia i lokalnych zmian koncentracji jonów występują w tkankach poddanych działaniu prądu stałego zjawiska polaryzacji elektrolitycznej, które zmniejszają wartość natężenia prądu.

Zjawisko to nie zachodzi dla prądów zmiennych i dlatego opór elektryczny tkanek dla prądu zmiennego jest znacznie mniejszy. Prąd stały używa się w lecnictwie do wprowadzania do organizmu różnego rodzaju jonów, do galwanizacji stabilnej i labilnej oraz kąpieeli komorowych.

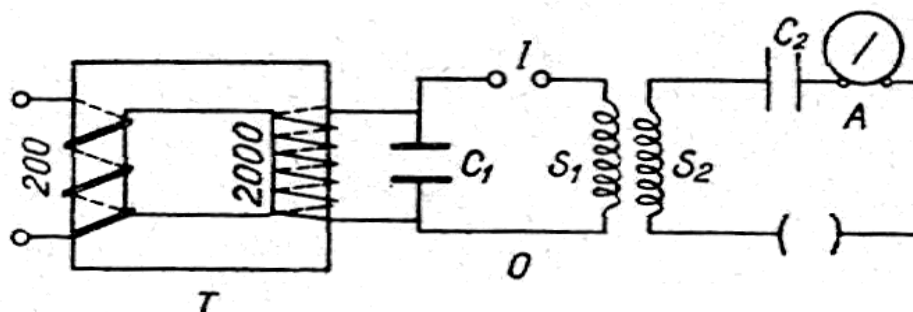
Aparat diatermiczny

Przenikanie prądów szybkozmiennych o dużym natężeniu, rzędu kilku amperów, przez organizm ludzki bez szkody dla niego zostało wykorzystane dla celów leczniczych. Mianowicie, jak wiemy, z przechodzeniem prądu elektrycznego związane jest wydzielanie się energii cieplnej w postaci ciepła Joule'a w ilości proporcjonalnej do kwadratu natężenia prądu. Przy prądach stałych i prądach zmiennych małej częstotliwości wydzielanie większej ilości energii cieplnej jest związane z przepływem prądu, który wywołuje w organizmie duże zmiany elektrochemiczne. Ponieważ prądy szybkozmiennne nie powodują takich dużych zmian, można je stosować do ogrzewania narządów zewnętrznych i wewnętrznych ciała. Ta metoda leczenia nosi nazwę **diatermii**.

Diatermia pozwala na ogrzewanie poszczególnych części wewnętrznych organizmu, podczas gdy ogrzewanie ciepłem z zewnątrz może sięgać zaledwie na parę milimetrów pod skórę, głównie ze względu na szybkie rozprowadzenie ciepła przez obieg krwi.

Zabiegi diatermiczne wywołują przejściowy spadek ciśnienia krwi w tętnicach, rozszerzanie się naczyń krwionośnych itd. Stosuje się je w reumatyzmie, podagrze, zapaleniach stawów, chorobach ginekologicznych itd.

Zasadniczy schemat aparatu diatermicznego pokazany jest na ryc.3. Prąd zmienny z sieci o napięciu około 200 woltów wchodzi do pierwotnego uzwojenia transformatora T i wychodzi przetworzony na prąd 2000 woltów. Zasila on obwód oscylatora O o pojemności C_1 , iskierniku I i cewce S_1 , w którym to obwodzie powstają prądy wielkiej częstotliwości. Z obwodem oscylatora sprzężony jest indukcyjnie drugi obwód drgań, zawierający pojemność C_2 cewkę S_2 , amperomierz ciepły A i elektrody, pomiędzy którymi znajduje się ciało pacjenta.



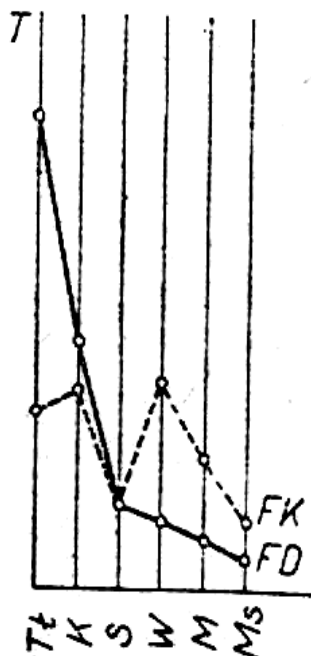
Ryc.3. Schemat diatermii.

Rozróżniamy obecnie dwa rodzaje prądów diatermicznych, zależnie od ich częstotliwości drgań; diatermię długofalową (długość fali około 300 metrów, tzn. częstotliwość drgań około miliona razy na sekundę) i diatermię krótkofalową (długość fali od 3 do 12 metrów, tzn. częstotliwość drgań od 25 do 100 milionów razy na sekundę)

Obydwa rodzaje prądów diatermicznych działają bardzo różnie na poszczególne tkanki ciała. Na przykład diatermia długofalowa ogrzewa przede wszystkim tkanki znajdujące się tuż pod

skórą, później mięśnie, a najmniej kości; natomiast diatermia krótkofalowa ogrzewa tkanki bardziej równomiernie, ale najintensywniej kości i wątrobę.

Ryc.4 przedstawia schematycznie krzywe rozkładu temperatury w różnych tkankach w zależności od rodzaju diatermii. Jak widzimy z ryciny, temperatury poszczególnych tkanek różnią się minimalnie w przypadku stosowania diatermii krótkofalowej (FK), a natomiast różnią się znacznie w przypadku stosowania diatermii długofalowej (FD), przy czym największą temperaturę osiąga tkanka tłuszczowa (Tł), mniejszą - tkanka kostna (K) i jeszcze mniejsze, ale już bliskie temperatury osiągają: skóra (S), wątroba (W), mózg (M) i mięśnie (Ms).



Ryc. 4. Krzywe rozkładu temperatury w tkankach przy diatermii długofalowej (FD) i krótkofalowej (FK).

W diatermii długofalowej przykładają się elektrody bezpośrednio do ciała pacjenta, natomiast w diatermii krótkofalowej elektrody znajdują się w pewnej odległości od ciała z tym, że musi się ono znajdować pomiędzy nimi w polu elektrycznym kondensatora.

Fulguracja, zimna kaustyka, elektrowstrząsy

Prądy wielkiej częstotliwości mogą jednak wywołać, przy zbyt dużym wydzielaniu ciepła albo przy wyładowaniach iskrowych, szkodliwe spalanie tkanek. W nowszych czasach wyładowania iskrowe z transformatorów wielkiej częstotliwości zostały wykorzystane do odkażania świeżych płaszczyzn cięcia przy operacjach (tzw. fulguracja), szczególnie gdy chodzi o zniszczenie ewentualnych zawiązków raka. Wytwarzany przy tym ozon ma również pewną rolę odkażającą.

Poza tym stosuje się prądy wielkiej częstotliwości w elektrochirurgii przy przecinaniu tkanek (tzw. elektrotomia). W pewnym miejscu ciała pacjenta przykładają się wtedy elektrodę o dużej powierzchni, a w miejscu cięcia przykładają się drugą elektrodę w postaci wąskiego i cienkiego noża platynowego (lancetu). Elektrody włączają się do źródła prądu wielkiej częstotliwości i wtedy ze względu na bardzo dużą gęstość prądu (rzędu 10^4 A/m^2), jaka wytwarza się na ostrzu lancetu, tkanki zostają na bardzo wąskiej przestrzeni po prostu spalone; powstające przy tym iskierki są wielkości mikroskopijnej, tak że przy szybkim ruchu noża ślad cięcia jest minimalny.

Jednocześnie następująca koagulacja przeciwdziała krwawieniu tkanek w miejscu cicia. Taki lancet elektryczny nazywamy kauterem wielkiej częstotliwości albo zimnym kauterem.

Do najnowszych działów elektromedycyny, stosowanych praktycznie dopiero od kilkunastu lat, należą: elektrowstrząsy i elektronarkoza. W obydwu przypadkach zjawisko polega na oddziaływaniu prądu elektrycznego na tkankę mózgową pacjenta. Różnica działania jest natury ilościowej.

Przy elektrowstrząsach przepuszczamy prąd zmienny (albo jednokierunkowy modulowany) przez głowę pacjenta w ciągu bardzo krótkiego czasu (od 0,1 s do 4 s), przy tym czas jest regulowany dokładnie w każdym aparacie przez specjalny wyłącznik automatyczny. Napięcie takiego prądu jest najczęściej rzędu stukilkudziesięciu woltów (w niektórych aparatach dochodzi do 300 woltów), natężenie prądu może się wahać w granicach od 150 do 800 miliamperów.

Przy elektronarkozie stosuje się napięcia mniejsze, tak żeby prąd elektryczny miał natężenie rzędu kilkudziesięciu do stu miliamperów, przy tym często można stosować również obok prądu zmiennego i prąd stały. Czas zabiegu jest przy tym dłuższy (rzędu kilku minut).

Aparaty do elektrowstrząsów i do elektronarkozy są zaopatrzone w przyrządy pomiarowe (amperomierze, woltomierze i omomierze) w urządzenia zabezpieczające pacjenta i obsługę przed zwiększeniem napięcia, przed oparzeniami i porażeniami w razie nagłego oderwania się elektrod od skóry pacjenta w czasie zabiegu.

Działanie mikrofal (fal radarowych i telewizyjnych)

Wraz z rozwojem radiotechniki, radiolokacji i teletechniki oraz licznymi zastosowaniami tych dziedzin w technice, w przemyśle i w wojskowości powstało bardzo poważne zagadnienie szkodliwego oddziaływania prądów wysokiej częstotliwości na żywy organizm. Szczególnie jest to o tyle ważne, że zaczęto wytwarzać źródła tych prądów o bardzo dużej mocy dochodzącej do 100 megawatów i o częstotliwości rzędu dziesiątków tysięcy megaherców. W związku z tym liczba ludzi znajdujących się w zasięgu oddziaływania takiego promieniowania elektromagnetycznego w. cz. i narażonych na jego oddziaływanie szybko wzrasta.

Obecnie ustalane są już normy dopuszczalnego maksymalnego dawkowania człowieka takim promieniowaniem w zależności od częstotliwości promieniowania od długości jednorazowego napromieniowania itd.

W przemyśle, w technice i w medycynie stosuje się promieniowanie w. cz. w granicach częstotliwości od 3 MHz (długość fali 100 m) do 3×10^5 MHz (długość fali 1 mm).

Pola elektromagnetyczne w. cz. w granicach od około 3 do 30 MHz stosuje się w przemyśle i w technice do termicznej obróbki metali i dielektryków, do suszenia, klejenia, spawania, polimeryzacji, sterylizacji żywności, niszczenia szkodników, do łączności radiowej itd.

Pola elektromagnetyczne o częstotliwości w granicach od 30 do 300 MHz stosuje się w łączności radiowej i telewizyjnej.

Specjalnie dużo uwagi poświęcono tzw. mikrofalom, których zakres długości fali wynosi 10^{-1} - 100 cm, a częstotliwość 3×10^2 - 3×10^5 MHz.

Do tego zakresu należą fale radarowe i fale telewizyjne. Mikrofałe graniczą z jednej strony z falami radiowymi ultrakrótkimi (o częstotliwości od 30 do 300 MHz), z drugiej strony z promieniowaniem podczerwonym (o częstotliwości od 3×10^6 do 4×10^8 MHz).

Cały zakres opisanego wyżej pola elektromagnetycznego w. cz. znajduje również zastosowanie w medycynie, między innymi w diatermii. Poza tym stanowi duży problem dla bezpieczeństwa i higieny pracy (bhp). Poza bezpośrednim działaniem fizjologicznym tych pól należy się jeszcze liczyć z działaniem promieniowania rentgenowskiego, jakie może być wytworzone ubocznie w generatorach mikrofalowych.

ELEKTRYCZNE WŁAŚCIWOŚCI KOMÓRKI

Właściwości błony komórkowej.

Komórki żywego organizmu są źródłem siły elektromotorycznej (SEM). Siła ta powstaje na skutek wędrówki jonów i zmiany ich koncentracji w różnych miejscach komórki. Takie przesunięcia jonów wywołują powstanie w tych miejscach różnicy potencjałów elektrycznych, która w pewnym momencie może wywołać impuls prądu elektrycznego przesuwający się wzdłuż tkanki i powodujący wyrównanie tych potencjałów

Z tego względu rozróżniamy w fizjologii potencjały spoczynkowe i tzw. potencjały czynnościowe towarzyszące impulsom prądu. Różnica potencjałów spoczynkowych powstaje najczęściej na skutek różnicy ruchliwości jonów przy ich przechodzeniu przez półprzepuszczalne błony komórek

Z tego względu obserwuje się na ogół różnicę potencjałów pomiędzy wewnętrzną częścią komórki i jej otoczeniem, a więc po dwóch stronach błony komórkowej. Błona taka stanowi zatem pewnego rodzaju kondensator elektryczny z warstwą dielektryka w środku, który może ulec naładowaniu do różnicy potencjałów rzędu od ułamka do kilkuset miliwoltów.

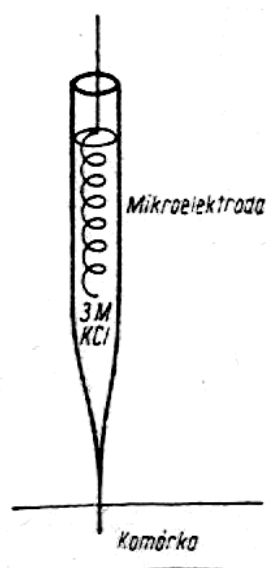
Jest to napięcie bardzo małe, ale ze względu na znikomo małą grubość błony komórkowej (rzędu od kilkudziesięciu do kilku tysięcy angstromów) mogą w niej powstawać pola elektryczne o natężeniu setek tysięcy woltów na centymetr.

Opór elektryczny błony komórkowej jest rzędu $10^7 \Omega/\text{m}^2$ lub $1,3 \times 10^{12} \Omega/\text{m}^3$, jej pojemność elektryczna wynosi około $1,8 \times 10^4 \mu\text{F}/\text{m}^2$. Opór elektryczny plazmy i płynów międzykomórkowych wynosi około $5 \times 10^7 \Omega/\text{m}^3$.

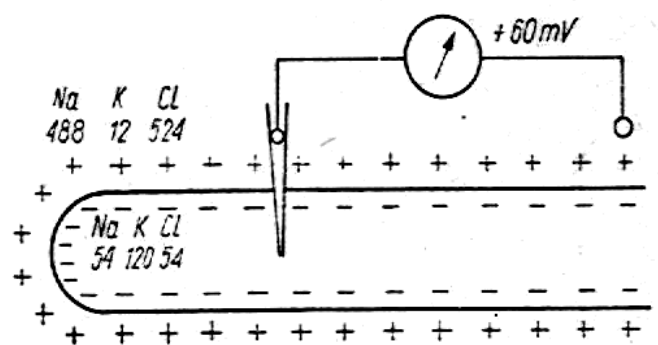
Taki stan błony nazywamy stanem spolaryzowanym. Jeżeli następuje rozładowanie tego kondensatora, mówimy, że nastąpiło pobudzenie komórki, a następnie jej depolaryzacja, tzn. wyrównanie potencjałów.

Badania doświadczalne potencjałów czynnościowych przeprowadza się za pomocą mikroelektrod. Mikroelektroda jest to rurka zakończona cienką kapilarą o średnicy rzędu jednego mikrometra (ryc.5), wypełniona 3 M KCl. W rurze znajduje się zwinięty spiralnie drucik srebrny pokryty AgCl_2 .

Mikroelektrodę wkuwa się za pomocą mikromanipulatora do wnętrza komórki, a drugą elektrodę, uziemioną umieszcza się na zewnątrz komórki w płynie fizjologicznym (ryc. 6).



Ryc. 5
Mikro-
elektroda.

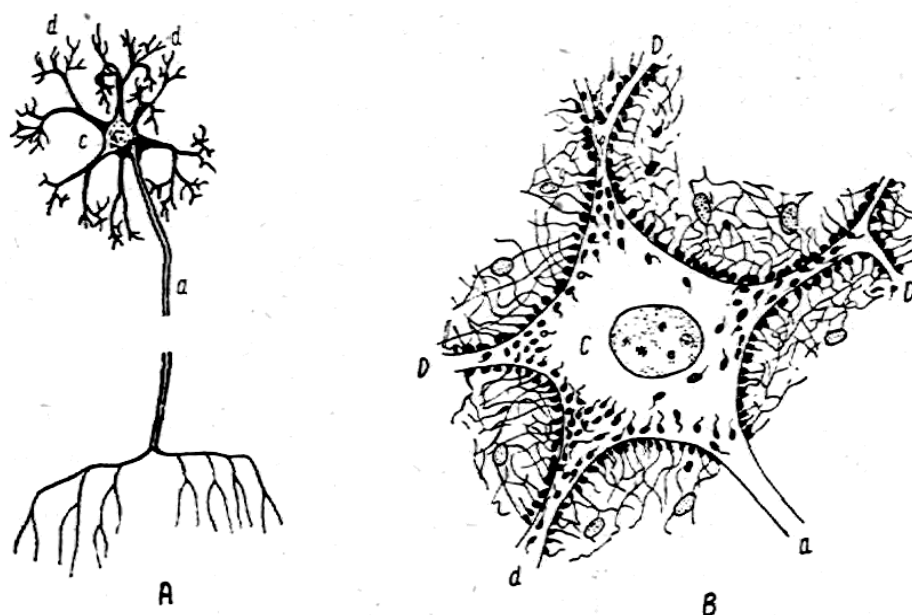


Ryc. 6 .Schemat rozmieszcze-
nia mikroelektrod w komórce
i rozkład stężeń jonów Na, K, Cl;
wg Karczewskiego.

Mikroelektrodę dołącza się do jakiegoś czułego miliwoltomierza, elektrometru czy oscyloskopu elektronowego. Dla zobrazowania procesów fizycznych zachodzących w żywej komórce musimy najpierw przypomnieć najważniejsze cechy fizyczne żywej komórki. Jako przykład weźmiemy komórkę nerwową, tzw. neuron; jest to jedna z najważniejszych komórek w organizmie; dokonuje ona nie tylko przemiany materii, ale również wytwarza impulsy elektryczne, transformuje je, przewodzi i dostarcza do innych komórek nerwowych i mięśniowych.

Neuron.

Neuron składa się z niewielkiej części centralnej (tzw. ciała neuronu) (ryc.7) i z wielu wypustek, tzw. d e n d r y t ó w, wychodzących z tej części na zewnątrz i rozgałęziających się; poza tym istotną częścią neuronu jest tzw. a k s o n, długa pojedyncza wypustka, zakończona również odgałęzieniami. Akson odgrywa rolę kabla elektrycznego, tzn. przewodnika elektrycznego owiniętego warstwami izolatora. Średnica aksona jest rzędu 1-20 mikronów, podczas gdy jego długość może dochodzić do rzędu metrów.



Ryc. 7 A — Komórka nerwowa (neuron): C — ciało neuronu; d — dendryty, a — akson; B — synapsy otaczające dendryty i ciało neuronu; wg Living Body.

Błona komórkowa otaczająca akson składa się z kilku warstw. Warstwa wewnętrzna, najbliższa środka, tzw. błona pobudliwa, jest zbudowana w postaci kondensatora; jego okładki są utworzone z dwóch warstw białka, a dielektrykiem jest warstewka tłuszczu, znajdująca się pomiędzy nimi. Grubość błony pobudliwej jest rzędu kilkudziesięciu angstromów. Na zewnątrz od tej błony znajdują się inne warstwy osłaniające akson (ryc. 8): osłonka mielinowa, osłonka nerwowa, osłonka Schwanna i tzw. endoneurium.

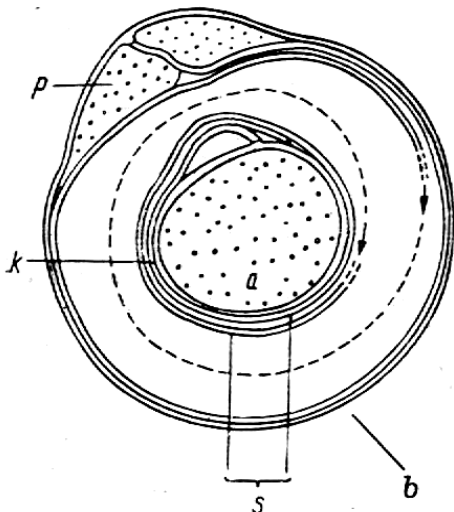
Ośłonka mielinowa składa się z krótkich odcinków włókna zbudowanego z białek i z fosfolipidów o długości od 0,5 do 3 mm, owijających akson. Zwoje osłonki mielinowej nie otaczają szczelnie aksonu na całej jego długości, ale tworzą pewne przerwy, które noszą nazwę przewężeń Ranviera (ryc 9); w tych przerwach akson jest odsłonięty. Na ryc. 10 pokazany jest szczegółowy schemat przewężenia Ranviera.

Na końcu aksonu znajdują się rozgałęzienia również nie osłonięte osłonką mielinową i zakończone zgrubieniami, wewnątrz których znajdują się cząstki acetylocholinu o średnicy rzędu kilkuset angstromów.

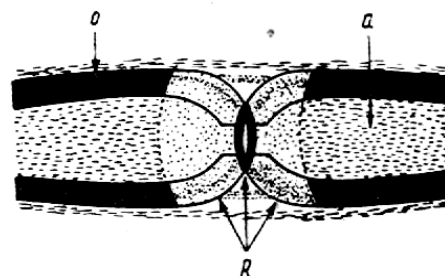
Rozkład potencjałów w komórce nerwowej

Badania wielkości i rozkładu potencjałów w komórkach wykazują, że na ogół wewnątrz komórki ma potencjał ujemny w stosunku do potencjału płynu pozakomórkowego, przy tym wartość tego potencjału w komórkach nerwowych jest rzędu 90 mV, w innych komórkach jest przeważnie mniejsza.

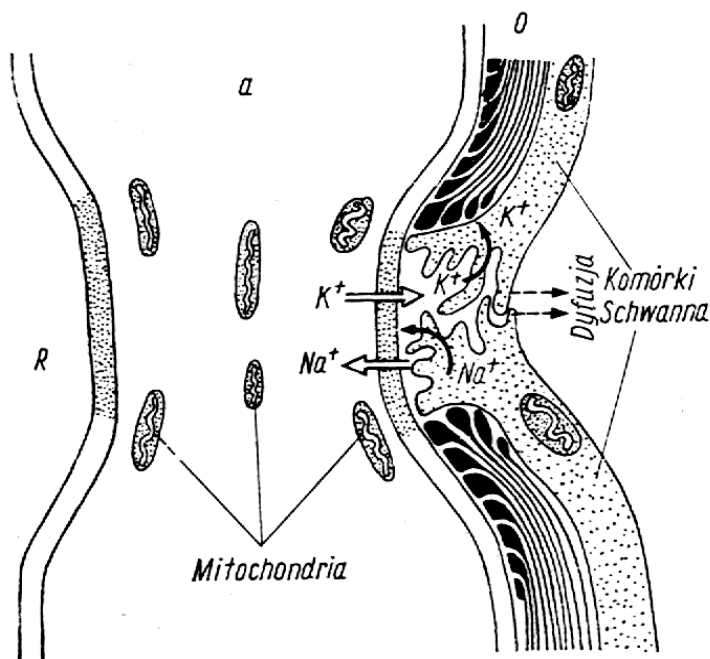
Poza tym wykazały one, że przyczyną powstawania tej różnicy potencjałów jest głównie różnica w stężeniach jonów sodu, potasu i chloru po obu stronach błony komórkowej.



Ryc. 8 Przekrój poprzeczny aksonu (a) i osłonki mielinowej; p — przestrzeń dla komórki Schwanna, k — aksolemma, S — osłonka Schwanna, b — błona komórki.



Ryc. 9 Akson z przewężeniem Ranviera: a — akson, o — osłonka mielinowa, R — przewężenie.



Ryc. 10 Szczegółowy schemat przewężenia Ranviera; wg Muralta.

Mianowicie, wewnątrz komórki stężenie jonów sodu (Na^+) jest około 10 razy niższe niż na zewnątrz, stężenie jonów chloru (Cl^-) jest od 10 do 65 razy niższe, a stężenie jonów potasu (K^+) jest około 10 do 50 razy wyższe niż na zewnątrz.

Wartości te wahają się w dosyć dużych granicach i są uzależnione od przepuszczalności błony komórkowej dla różnych jonów, przy tym półprzepuszczalność błony może być uwarunkowana różnymi czynnikami, np. duża ilość białka w komórce wpływa na powiększenie różnicy w koncentracji jonów.

Poza tym błony komórkowe różnych komórek wykazują niekiedy bardzo selektywne

właściwości przepuszczania różnych rodzajów jonów, np. powłoki czerwonych krwinek (erytrocytów) przepuszczają głównie aniony, membrany komórek nerwowych przepuszczają głównie potas, a nie przepuszczają sodu itd.

W tabeli poniżej podany jest przykład stężenia jonów w aksoplaźmie włókna Loligo i w płynie pozakomórkowym (w materiale świeżo spreparowanym).

Stężenie jonów K⁺, Na⁺ i Cl⁻ [w mol/kg] wewnątrz (a) i na zewnątrz (b) komórki

	a	b	Stosunek stężeń
K ⁺	410	10	41
Na ⁺	49	460	0
Cl ⁻	40	540	0.074

Różnica w stężeniach jonów po obydwu stronach błony komórkowej powoduje powstawanie w tych miejscach różnicy potencjałów elektrycznych oraz przeciwdziałające im działanie siły dyfuzji, która dąży do wyrównania stężeń. Różnica potencjałów, która równoważy siłę dyfuzji w stanie równowagi dynamicznej przy danej różnicy stężeń nosi w fizjologii nazwę potencjału równowagi dla danego jonu. Potencjał równowagi (E) określa się z równania Nernsta:

$$E = - \frac{RT}{FZ} \ln \frac{C_z}{C_w} \quad [11]$$

gdzie T oznacza temperaturę w skali bezwzględnej, F - stałą Faradaya, Z - wartościowość danego jonu, C_z i C_w - stężenia jonów danego rodzaju na zewnątrz i wewnątrz komórki.

Dla jonów jednowartościowych (Z = 1) w temperaturze ciała ludzkiego (t = 37°C, T = 310°K) wzór [11] przybiera formę prostszą

$$E = 61,5 \ln \frac{C_w}{C_z} \quad [12]$$

co np. dla potasu C_w/C_z = 41 daje wartość 95 mV bardzo bliską wartości doświadczalnej.

Opisany wyżej prosty mechanizm wytwarzania potencjałów elektrycznych w komórce nasuwa jednak jeszcze wiele niejasności i wątpliwości. Z tego względu rozpatrywano również możliwości szeregu innych procesów, które mogą tłumaczyć zjawiska elektryczne obserwowane w żywym organizmie. Pod uwagę brano tutaj niektóre procesy znane z innych działów fizyki, np. wpływ przenoszenia swobodnych elektronów, swobodnych protonów, mechanizm przenoszenia nośników prądu w półprzewodnikach elektronowych czystych i domieszkowych, właściwości elektryczne warstw podwójnych itd.

Na przykład Davies zwrócił uwagę, że, jak wiadomo, membrana komórkowa ma charakter lipoproteinowy, przy czym warstwa lipidów styka się z wodnymi roztworami wewnątrz i na zewnątrz komórki. Dlatego przy rozpatrywaniu elektrycznych potencjałów komórkowych biofizyk powinien znać właściwości takich powierzchni rozdziału.

Energia jonów na granicy rozdziału dwóch faz zależy od położenia (orientacji) dipoli w obu fazach. W wodzie energia jonów jest stosunkowo duża, w lipidach mała. Między fazami lipidową i wodną nie ma ostrego przejścia, istnieje strefa przejściowa. Dla przejścia drobiny z fazy wodnej w lipidową trzeba pokonać energię hydratacji. Przy przejściu z fazy lipidowej w wodną energia solwatacji jest dużo mniejsza (energia przejścia odpowiada długości odcinka). Ta różnica energii prowadzi do powstawania gradientów koncentracji i określa wielkość

współczynnika rozdziálu (rozkładu).

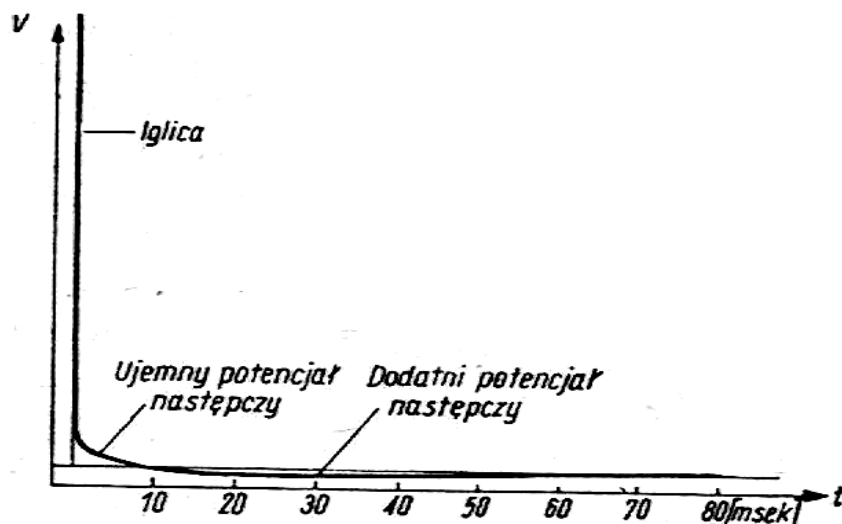
Teoretycznie różnica koncentracji, pojawiająca się w następstwie przenoszenia energii jonów poprzez membranę komórkową, może prowadzić do powstawania potencjałów elektrycznych znacznych wielkości. Taka energetyczna bariera uwarunkowuje małą przenikliwość, duży elektryczny opór i mniejszą dyfuzję przez membranę. Wszystkie te właściwości dadzą się obserwować w membranach żywych komórek.

Istniała również teoria oparta na teorii kwasowo-zasadowych potencjałów, w której wykazano, że dużo bioelektrycznych zjawisk można wyjaśnić przenoszeniem protonów. Wymiary protonów są znacznie mniejsze od rozmiarów jonów w metalu, natomiast gęstość ładunków protonów jest o wiele większa od gęstości ładunku tych jonów. Proton może przenosić się od drobiny do drobiny, podobnie jak elektron. Przewodnictwo takiego typu obserwuje się w kwasie siarkowym, a także w kwasach i w innych substancjach znajdujących się w stanie stałym.

Przy koncentracjach reagujących substancji obserwowanych w żywych komórkach (tkankach), różnica potencjałów może być równa kilkaset mV. Można więc powiedzieć, że powstawanie potencjałów bioelektrycznych może być uwarunkowane procesami protonowo-chemicznymi. Istnieje jeszcze kilka innych teorii powstawania potencjałów bioelektrycznych, ale nie będziemy ich tu omawiali.

Przenoszenie bodźców elektrycznych w komórkach.

Bodźcami pobudzającymi komórkę mogą być procesy chemiczne, mechaniczne, elektryczne, świetlne lub cieplne. W badaniach doświadczalnych w biofizyce i w fizjologii stosuje się najczęściej bodźce elektryczne, wprowadzając do wnętrza komórki mikroelektrodę i doprowadzając do niej odpowiedni potencjał. Jeżeli ten potencjał jest ujemny, tzn. jeżeli mikroelektroda jest katodą, to przy pewnej wartości tego potencjału (tzw. wartości progowej) może nastąpić depolaryzacja błony komórkowej i występuje w komórce gwałtowny wzrost potencjału (tzw. potencjał czynnościowy) w postaci ostrego wierzchołka (ryc. 11) na krzywej rozkładu potencjału.

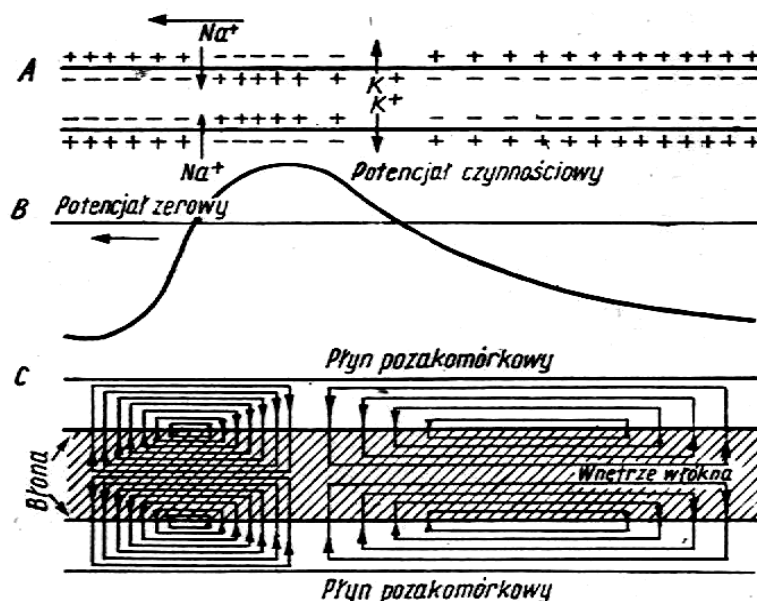


Ryc. 11 Rozkład potencjału czynnościowego; wg Gassera.

Odpowiada to jak gdyby "przebiciu elektrycznemu" błony komórkowej, gdyż jej opór elektryczny zmniejsza się w danym miejscu z wartości około $10^7 \Omega/\text{m}^2$ do wartości około $2,5 \times 10^5 \Omega/\text{m}^2$.

W miejscu "przebicia" przepuszczalność błony komórkowej dla jonów sodu staje się prawie całkowita, dzięki czemu następuje wyrównanie stężenia jonów sodu po obydwu stronach błony komórkowej, tzn. że jony sodu z zewnątrz "wlewają" się do wnętrza komórki niosąc ze sobą dodatni ładunek elektryczny, który znosi ujemny potencjał istniejący uprzednio wewnątrz komórki. W chwilę po tym rozpoczyna się inny proces, mianowicie gwałtowna dyfuzja jonów potasu z wnętrza komórki na zewnątrz, dzięki czemu zaczyna częściowo wzrastać różnica

potencjałów pomiędzy wnętrzem komórki i przestrzenią pozakomórkową. Obydwa procesy zachodzą zgodnie z działaniem sił gradientu elektrochemicznego, tzn. powodują rozładowanie ogniwa stężeniowego w komórce. Na ryc.12 pokazane są zmiany w czasie przepuszczalności błony komórkowej dla jonów sodu $p(\text{Na}^+)$ i dla jonów potasu $p(\text{K}^+)$ i wytworzony w ten sposób potencjał czynnościowy (V_c).

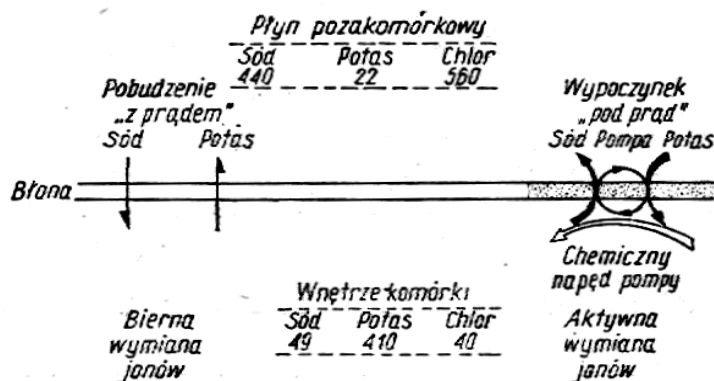


Ryc.12 Zmiany fizyczne i chemiczne zachodzące w komórce przy przewodzeniu potencjału czynnościowego: A — przesunięcie jonowe, B — rozkład potencjału wzdłuż włókna, C — kierunek prądów lokalnych; wg Karczewskiego.

Cały proces zachodzi w czasie rzędu paru milisekund. Ażeby przywrócić w komórce poprzedni stan typu ogniwa stężeniowego, należy sobie wyobrazić jakiś mechanizm, który przepycha jony sodu z powrotem na zewnątrz komórki, a jony potasu do jej wnętrza w obydwu przypadkach wbrew działaniom gradientów elektrochemicznych, tzn. powiększa różnice stężeń jonów każdego rodzaju po obydwu stronach błony komórkowej. Taki mechanizm nosi nazwę ogólnie pompy jonowej, a w szczególności pompy sodowej lub pompy potasowej (ryc. 13).

Jest on umieszczony całkowicie w błonie komórkowej, ale czerpie energię z protoplazmy. Energia ta powstaje z degradacji wysokoenergetycznych związków chemicznych w obecności tlenu, który jest konieczny dla działania tego mechanizmu.

Dalsze badania nad przebiegiem potencjału czynnościowego w czasie wykazały, że jest on funkcją bodźca i ma kształt pokazany na ryc. 14.

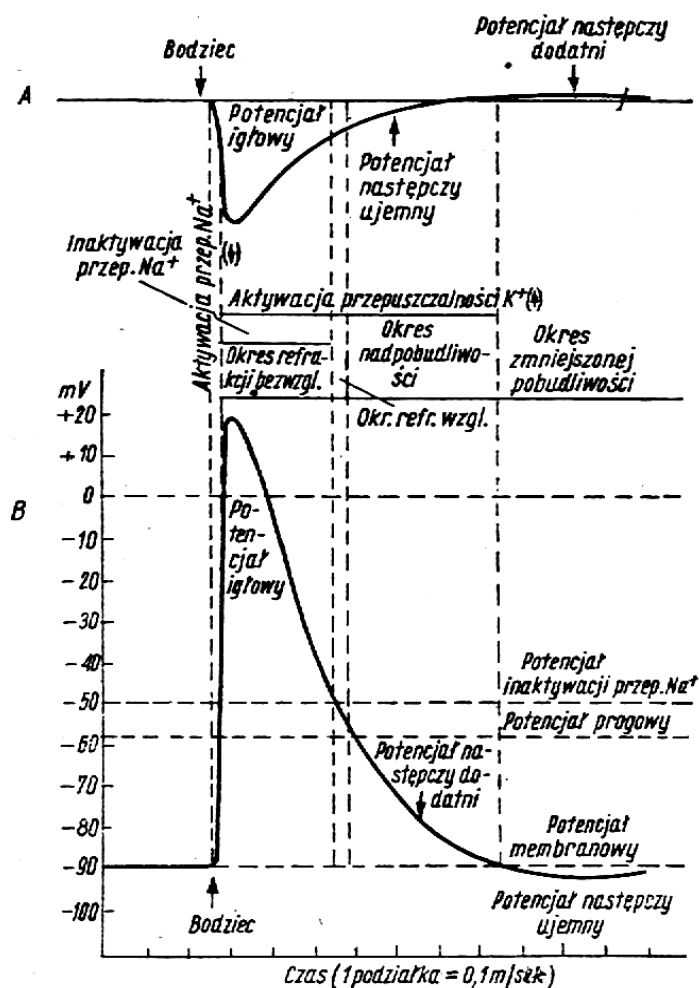


Ryc.13 Schemat działania pompy jonowej i przesunięć jonowych w czasie pobudzenia i w czasie spoczynku. Obszar zakreskowany odpowiada strefie działania pompy w czasie wypoczynku; wg Karczewskiego.

Po pierwszym, gwałtownym wzroście (w postaci iglicy) przy depolaryzacji błony komórkowej, następuje powolny spadek aż do wartości poniżej wartości początkowej i powrót

do stanu wyjściowego. Poszczególne odcinki krzywej potencjału czynnościowego noszą nazwę potencjału igłowego i potencjałów następnych: ujemnego i dodatniego.

Jak widać z powyższego opisu procesu pobudzania, mechanizm potencjału" czynnościowego jest bardzo skomplikowany, a trzeba jeszcze podkreślić, że w dużej mierze jeszcze hipotetyczny i niedostatecznie ugruntowany doświadczalnie pomimo bardzo wielu prac. prowadzonych przez wielu autorów różnymi metodami współczesnej fizyki, chemii i fizjologii. Część prac nad mechanizmem wędrówki jonów przez błonę komórkową wykonano za pomocą metod radioizotopowych stosując promieniotwórcze izotopy sodu i potasu. Powstanie potencjału czynnościowego jest tylko pierwszą częścią procesu przenoszenia impulsów pobudzenia przez komórki. Zmiany potencjału wewnątrz komórki są źródłem siły elektromotorycznej prądu elektrycznego, który może przepływać przez nerwy i włókna mięśni pomiędzy miejscami o różnym potencjale. W zależności od rodzaju przewodnika przewodzenie może być ciągłe albo przerywane (skokowe).

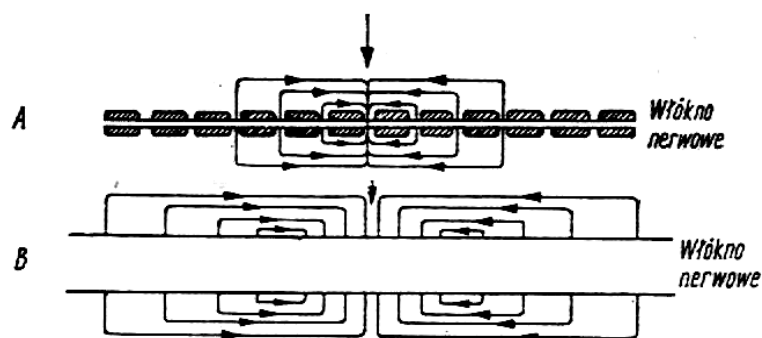


Ryc. 14 Zmiana potencjału czynnościowego w czasie. A — jednobiegunowe od-prowadzenie zewnątrzkomórkowe, B — potencjał wewnątrzkomórkowy; wg Houssaya.

Przewodzenie ciągłe powstaje w bezrdzennych włóknach, nerwowych, we włóknach mięśni szkieletowych i gładkich oraz w komórkach mięśnia sercowego. Przewodzenie skokowe odbywa się we włóknach mielinowych, w których występują przewężenia Ranviera. Szybkość przewodzenia skokowego jest większa niż szybkość przewodzenia ciągłego. Proces wytworzenia potencjału czynnościowego pod wpływem bodźca zewnętrznego i przekazania go dalej wzdłuż nerwu czy włókna mięśniowego zależy nie tylko od wielkości (siły) tego bodźca, ale i od czasu jego trwania i od poprzedniego stanu komórki.

W fizjologii przyjmuje się na ogół prawo "wszystko albo nic", tzn. prawo, które mówi, że dany bodziec działa w sposób warunkowy. Jeżeli jest on dostatecznie duży w stosunku do progu pobudliwości komórki, to wywoła impuls, jeśli niniejszy od tego progu, to nie wywoła nic. Istnieją jednak przypadki, że bodźce o mniejszej wartości wywołują pewien stan podprogowy w komórce, który zmienia jej zachowanie w stosunku do następnego takiego bodźca. Drugim parametrem grającym rolę przy pobudzaniu komórki jest czas działania bodźca; im krótszy

czas działania, tym bodziec musi być silniejszy. Wreszcie trzecim parametrem jest uprzedni stan komórki. Komórka wyładowana (depolaryzowana) nie reaguje przez pewien czas na bodźce, a zatem wykazuje jak gdyby "pewien czas martwy". Jest on rzędu 0,2 do 2 milisekund. Taki stan komórki nosi w fizjologii nazwę refrakcji bezwzględnej. W przypadku, kiedy bodźce przychodzą za często, komórka może zareagować tylko częściowo albo też wymaga zwiększonej siły bodźców dla normalnego zareagowania. Taki sposób reagowania nosi nazwę refrakcji względnej. Mechanizm przewodzenia bodźców we włóknie bezrdzeniowym jest bardzo skomplikowany. Nie znajduje on wyraźnej analogii w znanych w fizyce mechanizmach przewodzenia prądów elektrycznych w metalach, w półprzewodnikach czy w elektrolitach. Niekiedy porównuje go się raczej do mechanizmu działania lontu, w którym, jak wiadomo, płomień przenoszony jest powoli z miejsca do miejsca, czerpiąc energię z nagromadzonego w danym miejscu materiału palnego.



Ryc.15 Schemat linii prądu płynącego pomiędzy pobudzonym i niepobudzonym odcinkiem włókna nerwowego. **A** — włókno rdzeniowe, **B** — włókno bezrdzeniowe; wg Fultona.

Na ryc. 15 pokazany jest schematycznie obraz hipotetycznego mechanizmu przewodzenia bodźca wzdłuż włókna rdzeniowego (A) i bezrdzeniowego (B). Po lewej stronie rysunku jest stan normalny włókna, w części środkowej następuje pobudzenie (zmiana potencjału czynnościowego), w części prawej następuje repolaryzacja i powrót do stanu normalnego.

Zatem przed każdym potencjałem czynnościowym wędruje we włóknie fala pobudzenia, która obniża próg pobudliwości następnej warstwy włókna. Tego rodzaju prędkość przewodzenia bodźców jest rzędu 0,5 do 2 m/s. Włókna bezrdzeniowe znajdują się w układzie wegetatywnym.

Inny mechanizm przewodzenia bodźców elektrycznych zachodzi we włóknach rdzeniowych. Jest to mechanizm typu skokowego. Jak wiemy przewodzenie w komórce nerwowej odbywa się głównie wzdłuż aksonu, który sam jest dobrym przewodnikiem elektryczności (jego opór właściwy wynosi około 1 Ω m), a jest osłonięty na przeważającej swojej części przez osłonę nielinową, która jest bardzo dobrym izolatorem (jej opór właściwy jest rzędu $8 \times 10^7 \Omega$ m. Osłonka mielinowa ma szereg przerw, w których akson jest prawie odsłonięty i ulega przewężeniu. Mechanizm przewodzenia bodźców wewnątrz osłoniętych części aksonu pomiędzy przewężeniami Ranviera jest podobny do opisanego wyżej mechanizmu przewodzenia w układach bezrdzeniowych, natomiast w przewężeniach Ranviera obwód prądu zamyka się ponad osłonką nielinową wytwarzając nowy potencjał czynnościowy w następnym przewężeniu. A więc potencjał czynnościowy jak gdyby przeskakuje z jednego przewężenia Ranviera do drugiego, przy tym szybkość takiego przenoszenia bodźców jest zależna od rodzaju włókien nerwowych.

Rozróżnia się trzy rodzaje włókien nerwowych o następujących właściwościach elektrycznych: grupę A z podgrupami α , β , γ o średnicy rzędu 1-20 μ , prędkości przewodzenia rzędu 100-150 m/s i o czasie martwym rzędu 0,5 m; grupę B o średnicy około 3 μ , o prędkości przewodzenia około 10 m i o czasie martwym rzędu 1,2 m; grupę C - najcieńsze, bezrdzeniowe, o prędkości przewodzenia około 1 m i o czasie martwym rzędu 2 ms.

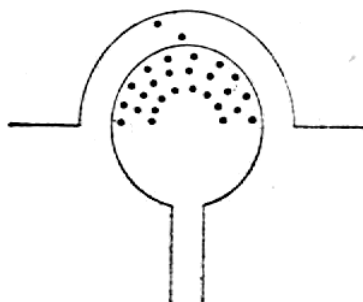
Nerw jest złożony z różnych grup wielu tysięcy poszczególnych włókien nerwowych, tak że bodziec elektryczny wędruje różnymi drogami, dając w sumie w miejscu doprowadzenia

złożony potencjał czynnościowy nerwu. Przekazywanie bodźców elektrycznych z jednego ośrodka do drugiego odbywa się najczęściej za pomocą dwóch lub więcej neuronów; zatem droga przewodzenia bodźca składa się z następujących elementów: dendryt, ciało neuronu, akson, zakończenie aksonu, tzw. synapsa, dendryt drugiego neuronu itd.

Synapsa jest elementem nie zupełnie jeszcze poznany. Znajduje się ona pomiędzy dwoma neuronami. Zakończenie jednego neuronu od początku drugiego neuronu dzieli przerwa o szerokości około 700 Å. W przerwie tej znajdują się czynne substancje chemiczne, takie jak acetylocholina, adrenalina i serotonina. Noszą one nazwę mediatorów. Poszczególne drobiny mediatora zostają uwolnione z tzw. pęcherzyków synaptycznych przez energię potencjału czynnościowego, przedostają się przez przerwę synaptyczną i działają na błonę osłaniającą drugi neuron. Działanie to wywołuje szereg drobnych miniaturowych potencjałów postsynaptycznych, które w sumie depolaryzują błonę i wytwarzają w niej potencjał czynnościowy.

Proces ten jest o tyle bardziej skomplikowany, że do błony jednego neuronu mogą dochodzić jednocześnie potencjały czynnościowe z tysięcy innych neuronów, przy tym niektóre z tych potencjałów mogą wywołać działanie przeciwne, hamujące pobudzenie. Synapsa wykonuje zatem szereg bardzo skomplikowanych czynności: segregujących, wybiórczych i sumujących bodźce i przekazujących je dalej, do następnego neuronu czy ostatecznie, do mięśnia.

W tym ostatnim przypadku bodziec trafia nie do kolejnej synapsy, ale do płytki końcowej, w której odbywają się procesy podobne jak w synapsie z udziałem odpowiednich mediatorów, tylko efekt końcowy jest inny, mianowicie skurcz mięśnia, a więc wykonanie pracy mechanicznej.



**Ryc. 16 Schemat modelu
płytki końcowej i przecho-
dzenie acetylocholiny przez
szczeliny.**

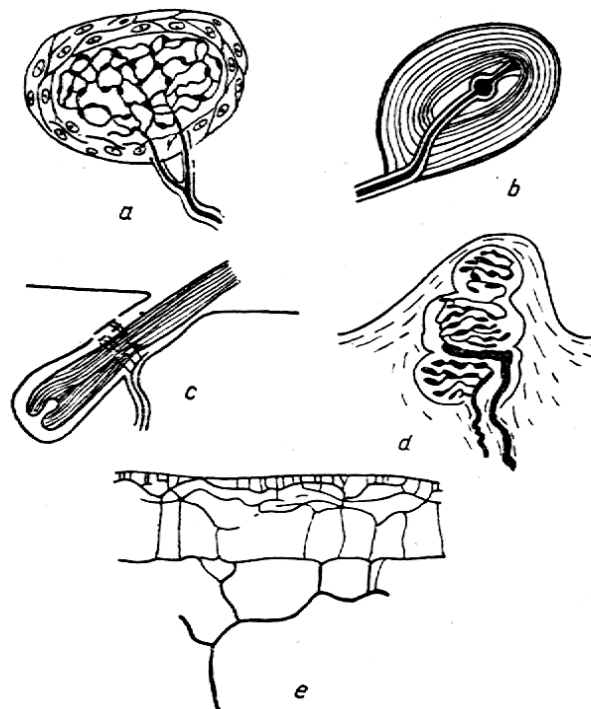
Każdy pojedynczy neuron ruchowy może pobudzić całą grupkę włókien mięśniowych, tzw. jednostkę ruchową. Wartość potencjału czynnościowego jednostki ruchowej jest stosunkowo mała, wynosi zaledwie około 0,5 mV, czas trwania tego potencjału wynosi około 12 milisekund. Droga przewodzenia bodźców elektrycznych przez układ neuron, synapsę, mięsień jest jednokierunkowa, tzn. że w drugą stronę bodźce w tym układzie nie przechodzą, głównie z uwagi na działanie synapsy. Ponieważ jednak pobudzenie mięśnia jest z reguły powiązane z wysłaniem bodźca czuciowego w drogę powrotną do ośrodkowego układu nerwowego, musi istnieć również osobna linia przewodzenia tego bodźca dla zasygnalizowania wykonania pierwotnego polecenia, ewentualnego skorygowania go czy wreszcie wykazania uczucia bólu.

W tym powrotnym układzie przewodzącym główną rolę na początku odgrywają tzw. receptory, specjalne komórki, czy zakończenia nerwowe wyczulone selektywnie na poszczególne rodzaje bodźców: mechanicznych, akustycznych, chemicznych, świetlnych, elektrycznych itd. (ryc.17).

Specjalną grupę/receptorów stanowią komórki przestrzegające przed uszkodzeniem czy zniszczeniem organizmu i sygnalizujące takie objawy przez wywoływanie bólu.

Istnieje bardzo wiele różnych rodzajów receptorów o różnej strukturze i różnych zadaniach funkcjonalnych (tzw. mechanoreceptory, chemoreceptory, termoreceptory itd.). W pobudzonym

receptorze pojawia się tzw. potencjał generujący, który wysyła bodźce elektryczne pojedyncze albo grupowe w zależności od rodzaju i wielkości pobudzenia. Częstotliwość wysyłania tych bodźców elektrycznych nerwu charakteryzuje rodzaj przesyłanej informacji.

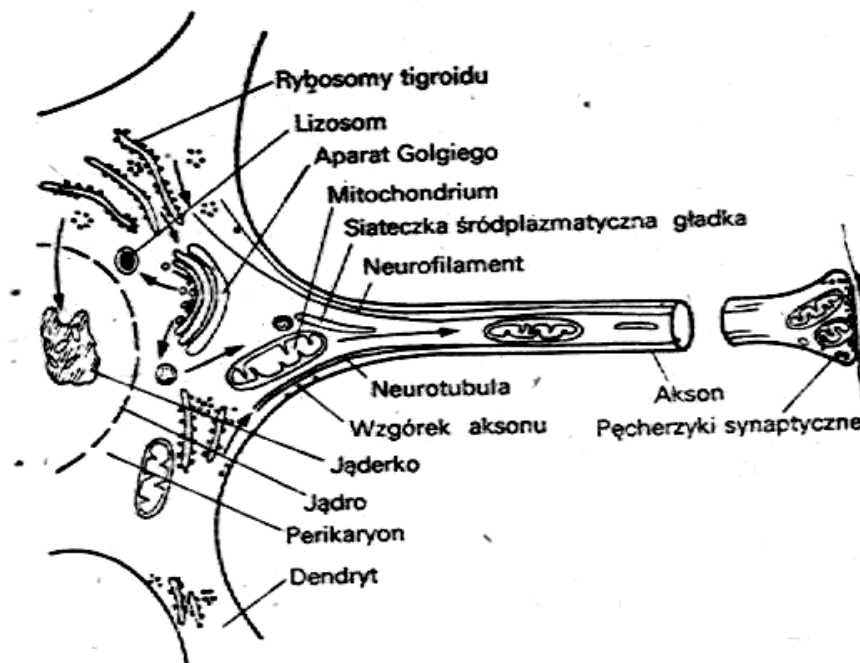


Ryc. 17 Różne receptory skóry: a — tarczki Merkela, b — ciała Krausego, c — receptory koszyrkowe mieszków włosowych, d — ciała Meissnera, e — wolne zakończenia włókien nerwowych; wg Houssaya.

Można zatem powiedzieć, że informacje pochodzące z receptorów są kodowane za pomocą modulacji częstotliwości. Jak wiemy z innych dziedzin nauki, taki sposób przekazywania informacji jest znacznie dokładniejszy niż przekazywanie informacji za pomocą modulacji amplitudy i podlega znacznie mniejszym zniekształceniom. Reasumując rozważania ostatnich paragrafów, można powiedzieć, że droga, jaką przebiegają potencjały czynnościowe w ustroju, rozpoczyna się od źródła pierwszej podniety, które znajduje się najczęściej w receptorze obwodowym, przebiega zawiłą drogę poprzez dendryty, ciało neuronu, akson, synapsę, inne neurony, mięsień i dochodzi poprzez receptory i neurony do mózgu. Cały ten proces trwa zaledwie ułamek sekundy.

BUDOWA KOMÓRKI NERWOWEJ JAKO WYRAZ PRZYSTOSOWANIA DO FUNKCJI

Jednostkami budulcowymi tkanki nerwowej są - podobnie jak to ma miejsce w tkankach innego rodzaju - komórki. Komórki nerwowe, zwane neuronami lub neurocytami, różnią się pod względem strukturalnym w zasadniczy sposób od komórek innego typu (ryc.18).



Ryc. 18 Schemat neuronu wraz ze strukturami wewnątrzkomórkowymi (zmodyfikowane wg Droza i Koeniga).

Ta odmienność strukturalna-polegająca przede wszystkim na istnieniu długich wypustek komórkowych -jest wyrazem przystosowania neuronu do jego funkcji. W neuronie wyróżnić można następujące cztery strefy czynnościowe związane z podstawowymi jego funkcjami: a) wejście, b) inicjacja impulsów, c) przewodzenie impulsów i d)

wyjście. Wymienione tu struktury neuronu mogą się wykształcić w rozmaity sposób, co jest podstawą klasyfikacji komórek nerwowych na różne typy strukturalni-czynnościowe.

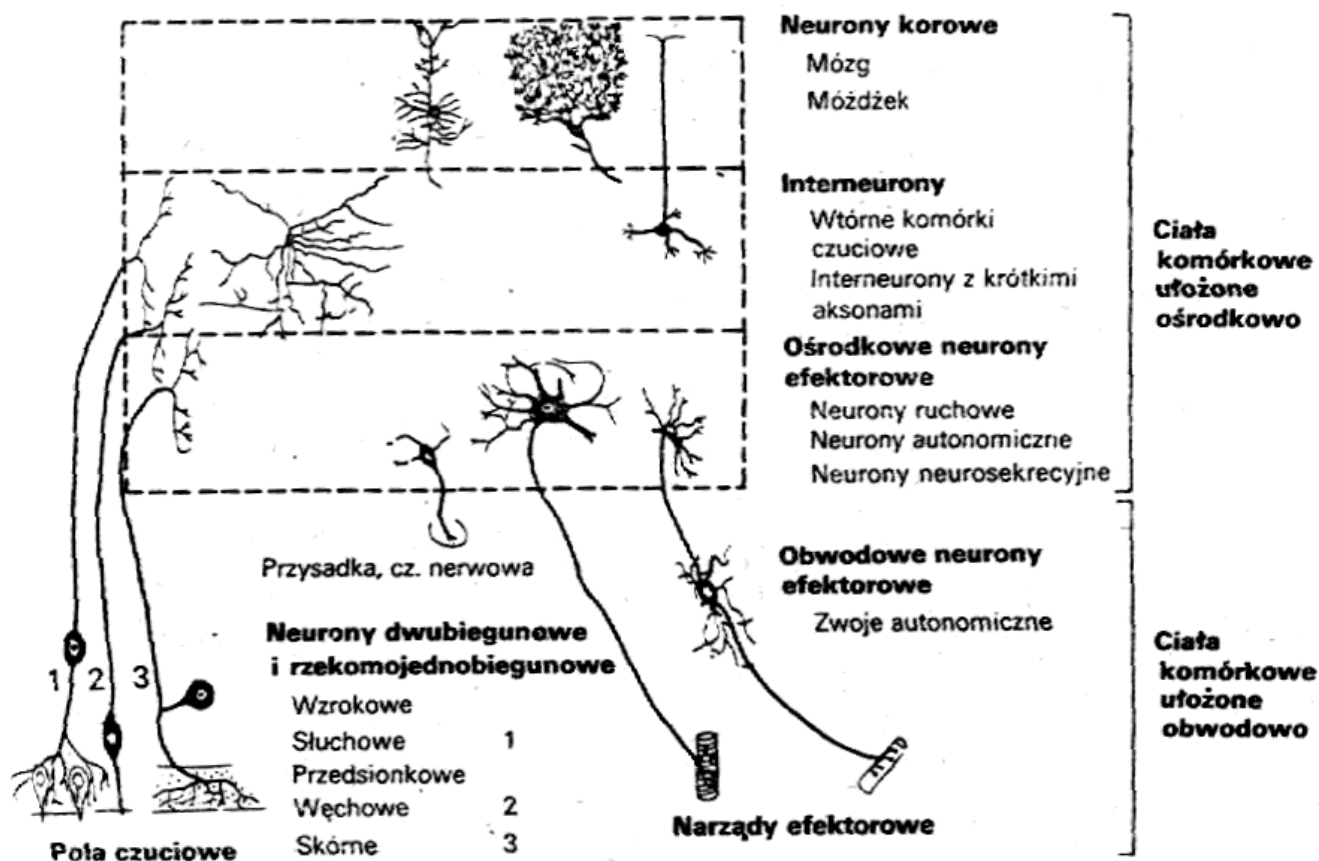
Podstawowe struktury neuronu

Mimo że zarówno w różnych obszarach układu nerwowego u tego samego osobnika, jak też u zwierząt różnych gatunków - spotkać można neurony o rozmaitych kształtach, podstawowe struktury wszystkich neuronów są w zasadzie jednakowe (ryc.18). Pierwszym podstawowym elementem strukturalnym komórki nerwowej jest ciało neuronu (soma). Ciała neuronów mogą być rozmaitego kształtu: od okrągłych poprzez gwiaździste do wrzecionowatych. Ciało neuronu przechodzi w dwa typy wypustek: aksony (dawna nazwa: wypustki osiowe) i dendryty (dawna nazwa: wypustki protoplazmatyczne),

Akson, czyli neuryt, jest głównym rejonem przewodzenia w neuronie, a więc główną drogą komunikowania się z inną komórką nerwową, lub też z komórką efektorą (np. mięsień, gruczoł). Ze względu na spełniane funkcje akson, który jest pojedynczą wypustką neuronu, zwykle charakteryzuje się dużą długością. Akson bierze swój początek w obszarze neuronu zwanym odcinkiem początkowym aksonu. Jest to element inicjacji impulsów. Zakończenie aksonu, zwykle mniej lub bardziej rozgałęzione, tworzy element wyjścia neuronu.

Ważnym z punktu widzenia czynności neuronu elementem strukturalnym są osłonki aksonu. W budowie struktur osłonowych uczestniczą komórki należące do tkanki glikowej. W

zależności od typów osłonek okrywających akson, wyróżniamy włókna nerwowe z osłonką mielinową., czyli rdzenne, i włókna nerwowe bez osłonki mielinowej, czyli bezrdzenne.



Ryc. 19 Rodzaje neuronów układu nerwowego ssaków (wg Ganonga).

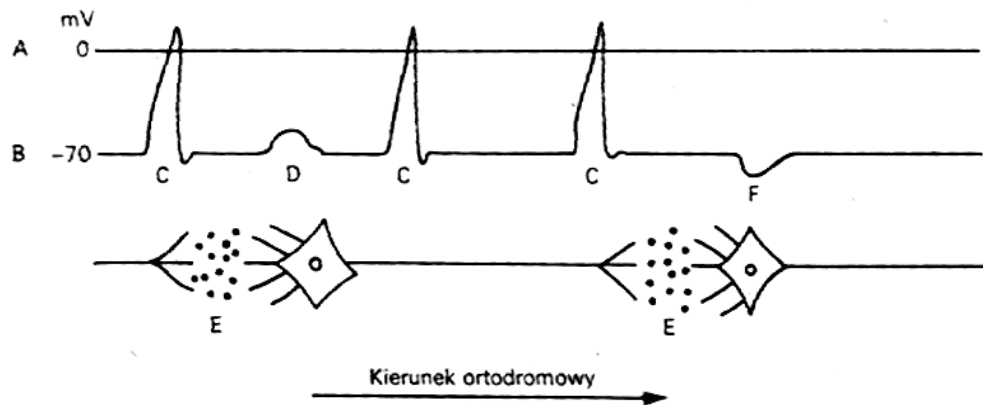
We włóknach rdzennych akson otoczony jest dwoma osłódkami: osłódką mielinową oraz neurolemą (dawna nazwa: osłódka Schwanna). Osłódka mielinową nie ma charakteru ciągłego. Przerywa się ona w regularnych odstępach, tworząc cieśni węzłów (dawna nazwa: węzły Ranyiera, przewężenia Ranviera). Neurolemą pokrywa osłódkę mielinową na przestrzeni całego włókna, a więc także w obrębie cieśni węzłów.

Drugim rodzajem wypustek neuronu są deadryty, stanowiące element wejścia neuronu. Występują one w neuronie w znacznej niekiedy liczbie i zwykle są krótsze od aksonu, a poza tym cechują się posiadaniem licznych drzewiastych rozgałęzień.

Spośród struktur specyficznych dla komórek nerwowych, które pojawiły

się w neuronach na skutek specjalizacji czynnościowej, należy wymienić: neurofibryle, tigroid (dawna nazwa: substancja Nissla) oraz wyspecjalizowaną w zakresie recepcji i przepuszczalności błonę komórkową. Neurofibryle są jednym z najbardziej charakterystycznych tworów występujących w komórce nerwowej. Są to cieniutkie nitkowate twory, wnikające do wszystkich wypustek neuronu. Jednostką mniejszą od neurofibryli są neurofila-menty. Struktury te tworzą układ neu-rotubuli. Niegdyś przypisywano neuro-tubulom zasadniczą i dominującą rolę w przenoszeniu pobudzenia wzdłuż aksonu. Dzisiaj nie ulega wątpliwości, że układ neurotubuli obok funkcji budulcowej odgrywa ważną rolę w transporcie neurohormonów, czy też przekazywaniu nerwowych z ciała neuronu do zakończeń aksonu. Grudki substancji zasadochłonnej tworzące tigroid są rozsiane w całej cytoplazmie neuronu. Brak ich w aksonie; w dendrytach występują obficie. Tigroid, składający się z rybosomów i siateczki śródplazmatycznej, jest skupiskiem kwasów nukleinowych i odgrywa ważną rolę w procesach przemiany materii i energii w neuronie.

POTENCJAŁY ELEKTRYCZNE KOMÓRKI NERWOWEJ



Ryc. 20 Potencjały bioelektryczne wnętrza komórki nerwowej. A — poziom zerowy, B — poziom potencjału spoczynkowego, C — potencjały iglicowe, D — postsynaptyczny potencjał pobudzający, E — synapsa chemiczna, F — postsynaptyczny potencjał hamujący.

Jednym z podstawowych przejawów życiowych komórki nerwowej jest jej zdolność do generowania i przewodzenia potencjałów elektrycznych. Te właściwości komórki nerwowej zlokalizowane są przede wszystkim w błonie komórkowej, która oddziela neuron od jego otoczenia.

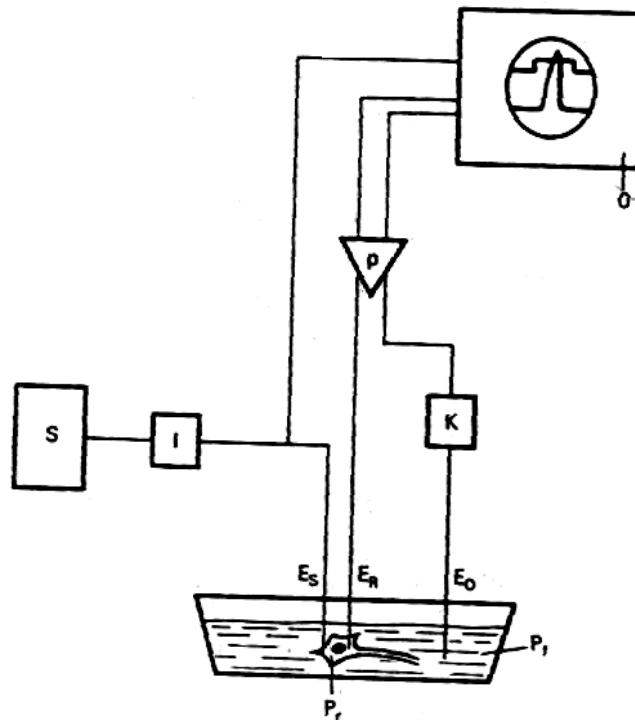
Wśród potencjałów bioelektrycznych można wyróżnić potencjał spoczynkowy oraz potencjały czynnościowe. Potencjał spoczynkowy występuje, gdy neuron nie jest pobudzany. Potencjały czynnościowe towarzyszą stanowi czynnemu neuronu (ryc.20).

Metody badania potencjałów bioelektrycznych

Istotny postęp w badaniach, neurofizjologicznych szedł zawsze w parze z rozwojem technik badawczych. Wśród technik stosowanych, do badania neuronu na pierwsze miejsce wysuwają się metody rejestracji potencjałów bioelektrycznych. Każda z takich metod zawiera w sobie dwa zasadnicze elementy: element stykający się bezpośrednio ze źródłem generowania potencjałów bioelektrycznych (elektroda) oraz część rejestrującą (np. oscyloskop). Między te dwa podstawowe elementy wyposażenia rejestrującego wprzęgnięte są zazwyczaj dodatkowe urządzenia w postaci np. przedwzmacniaczy.

W związku z tym, że neurony generujące potencjały bioelektryczne wytwarzają wokół siebie pole elektryczne, poza rejestracją z wnętrza komórki istnieje także możliwość rejestracji zewnątrzkomórkowej potencjałów bioelektrycznych.

Doświadczenia z wewnątrzkomórkowym pomiarem potencjałów bioelektrycznych w neuronie datują się od lat 1939-1940, kiedy to czterem uczonym (Cole, Curtis, Hodgkin, Huxley) udało się po raz pierwszy wprowadzić mikro-elektrodę do olbrzymiego aksonu mątwy, zwanego tak ze względu na średnicę dochodzącą do 1 mm. Podstawowe elementy techniki wewnątrzkomórkowej rejestracji potencjałów bioelektrycznych neuronu przedstawiono na ryc.21

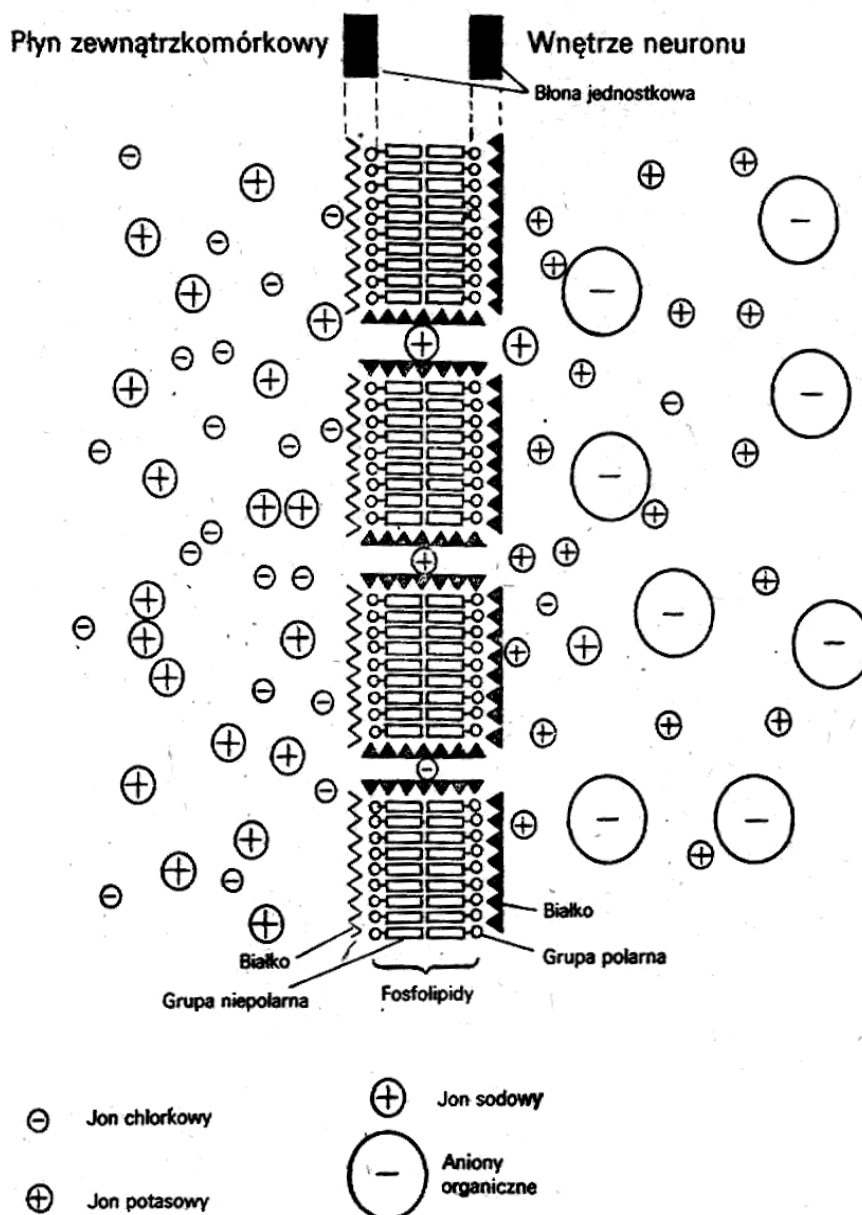


Ryc. 21. Metoda wewnątrzkomórkowego badania potencjałów bioelektrycznych neuronu. E_R — elektroda rejestrująca, E_s — elektroda stymulująca, E_o — elektroda obojętna, P_t — roztwór fizjologiczny soli, P_r — preparat, K — kalibrator, p — przedwzmacniacz, O — oscyloskop, I — jednostka izolująca, S — stymulator.

Obok metod rejestracji potencjałów bioelektrycznych ważnym elementem techniki badania komórki nerwowej jest metoda stymulacji. W zakresie badań nad pojedynczą komórką nerwową stosuje się obecnie przede wszystkim

stymulacją elektryczną i stymulację chemiczną. Metoda dokomórkowego wprowadzania substancji chemicznych nosi nazwę mikroelektroforezy lub jontoforezy.

Potencjał spoczynkowy i jego geneza



Ryc. 22 Hipotetyczna struktura błony neuronu (zmodyfikowane wg Usherwooda). Rozmieszczenie jonów po obydwu stronach błony przedstawiono schematycznie, a ich względne średnice w odniesieniu do jonów uwodnionych.

Powierzchnia nieuszkodzonego ciała neuronu, czy też jego włókien w spoczynku jest izopotencjalna. Znaczy to, że pomiędzy dwoma dowolnymi punktami na powierzchni neuronu nie ma różnicy potencjałów. Stwierdzana w neuronach w czasie spoczynku różnica potencjałów między wnętrzem a otoczeniem nosi nazwę potencjału spoczynkowego lub błonowego potencjału spoczynkowego. Różnica ta wynosi kilkadziesiąt miliwoltów.

Poza wieloma podobieństwami w składzie cytoplazmy neuronu i jego płynnego środowiska zewnętrznego istnieją istotne różnice w dystrybucji jonów w obydwu tych środowiskach (ryc. 22). W środowisku zewnątrzkomórkowym głównym anionem jest chlor, którego stężenie jest tu około 10 razy wyższe niż w cytoplazmie neuronu. U źródeł tego nierównomiernego rozłożenia jonów chlorkowych leży przede wszystkim obecność w cytoplazmie dużych, nieprzechodzących przez błonę anionów organicznych. W związku z tym można powiedzieć, że rozmieszczenie jonów chlorkowych w środowisku wewnątrz- i zewnątrzkomórkowym jest wynikiem równowagi Donnana.

Z jonów odgrywających pierwszorzędową rolę w genezie zjawisk bioelektrycznych należy wymienić sód i potas. Sód jest około 10 razy bardziej stężony na zewnątrz neuronu niż w jego wnętrzu. Stężenie natomiast potasu jest 30 razy wyższe w cytoplazmie neuronu niż w jego otoczeniu. Jeśli chodzi o sód, to biorąc pod uwagę ujemność wnętrza neuronu, łatwo dojść do

wniosku, że siły wynikające zarówno z gradientu stężeń, jak i gradientu elektrycznego działają w tym samym kierunku (dokomórkowo). Obserwowane asymetrie w dystrybucji jonów sodowych i potasowych w środowisku zewnątrz i wewnątrzneuronnym nie mogą być wynikiem biernych procesów błonowych. Asymetrie te są wynikiem aktywnego transportu jonów sodowych i potasowych. Transport ten dochodzi do skutku dzięki działalności w błonie neuronu systemu pomp sodowych i potasowych.

U podstaw tych układów leży działalność specjalnego błonowego enzymu transportującego - adenozynotrójfos-fatazy (Na-K-ATP-aza), aktywowanej przez Na^+ i K^+ . Proces transportu jonów przebiega ze zużyciem adenozyno-trójfosforanu (ATP). W następstwie hydrolizy ATP powstaje adenozynodwu-fosforan (ADP) i fosforan nieorganiczny. W czasie powyższej reakcji wyzwala się energia pochodząca z hydrolizy bogatego energetycznie wiązania fosforanowego. Energia ta zużywana jest do transportu jonów wbrew gradientowi stężeń. Szereg cech czynnościowych Na-K-ATP-azy pokrywa się z cechami pompy jonowej, wskazując na ścisły związek tych dwóch mechanizmów.

Częściowe lub całkowite zahamowanie działania pompy sodowo-potasowej nastąpić może poprzez: 1) niedotlenienie spowodowane brakiem tlenu lub działaniem inhibitorów oddychania komórkowego, 2) zastosowanie wybiórczych inhibitorów pompy, np. ouabainy i pokrewnych glikozydów" nasyconych, 3) spadek temperatury, powodujący obniżenie metabolizmu komórkowego.

Na skutek aktywności transportującej ATP-azy jony sodowe nieustannie usuwane są na zewnątrz komórki, a do jej wnętrza wprowadzane są jony potasowe, przy czym stosunek wymiany wynosi dwa jony potasowe za trzy jony sodowe. Transportująca ATP-aza aktywowana jest przez jony sodowe działające na wewnętrzną powierzchnię komórki. Im większy jest wobec tego napływ tych jonów do wnętrza na skutek biernej dyfuzji, tym większa jest aktywność pompy jonowej i tym większy czynny transport na zewnątrz. Na skutek istnienia tego ujemnego sprzężenia zwrotnego czynny transport sodu jest równy biernej dyfuzji tych jonów, wobec czego średni ich przepływ przez błonę pozostającą w stanie spoczynku równa się 0.

W poprzek błony neuronu istnieją wyraźne gradienty stężeń dla jonów sodowych, potasowych i chlorkowych (ryc. 23). Pozostająca w spoczynku błona neuronu jest wysoce przepuszczalna dla jonów potasowych, mniej dla chlorkowych i minimalnie dla jonów sodowych. Jeżeli przepuszczalność błony neuronu dla jonów potasowych przyjmujemy jako jedność, to przepuszczalność dla omawianych trzech jonów (K^+ , Cl^- , Na^+) wynosi odpowiednio-1 : 0,45 : 0,04. Biorąc to pod uwagę, a także pamiętając, że istnieje asymetria w stężeniach jonów w poprzek błony, łatwo dojść do wniosku, że w poprzek błony neuronu musi istnieć różnica potencjałów. Te różnice potencjałów określamy, jako potencjał spoczynkowy. Wielkość błonowego potencjału spoczynkowego określana jest głównie przez wielkość gradientu stężeń dla jonów potasowych i zbliża się do potencjału równowagi dla potasu.

Tak rozumianą spoczynkową różnicę potencjałów wyraża się wzorem:

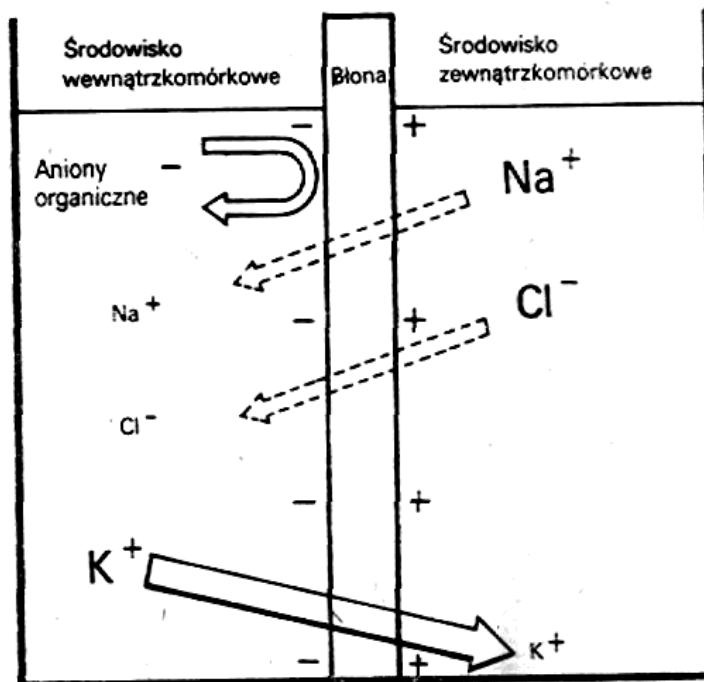
$$E_K = \frac{RT}{F} \ln \frac{[\text{K}^+]_z}{[\text{K}^+]_w} \quad [13]$$

gdzie:

E_K — potencjał równowagi dla jonów K^+ ,
 R — stała gazowa, T — temperatura absolutna, F — stała Faradaya, $[\text{K}^+]_z$ $[\text{K}^+]_w$ —
 stężenia jonów potasowych na zewnątrz
 i wewnątrz neuronu.

Z wzoru [13] łatwo wyciągnąć wniosek, że zwiększenie stężenia potasu w środowisku zewnątrz neuronu prowadzić musi do zmniejszania różnicy potencjału w poprzek błony neuronu. Należy pamiętać, że w omawianym przypadku mamy do czynienia z neuronem zachowującym się jak elektroda potasowa; nie można jednak wykluczyć udziału innych jonów w kształtowaniu spoczynkowej różnicy potencjałów. Ich rola jest jednak minimalna.

Ponieważ prawidłowy przebieg potencjałów czynnościowych i wszystkich. zjawisk z nimi związanych, a przede wszystkim przewodnictwa stanu czynnego, zależy od potencjału spoczynkowego, na tle którego komórki zostały aktywowane, prawidłowa czynność komórek



Ryc. 23 Parametry elektrochemiczne w neuronie w czasie spoczynku (zmodyfikowane wg Usherwooda). Lewy przedział zawiera wysokie stężenie K^+ i naładowane ujemnie cząsteczki organiczne (białka) oraz niskie stężenia Na^+ i Cl^- , podczas gdy w prawym przedziale rzecz ma się odwrotnie. Dyfuzja jonów potasowych bierze główny udział w tworzeniu potencjału spoczynkowego. Na rycinie nie przedstawiono pompy sodowo-potasowej.

zależy w dużym stopniu właśnie od prawidłowego potencjału spoczynkowego.

W stanach chorobowych przebiegających ze znacznego stopnia hiperpotasemii, tj. zwiększeniem stężenia jonów potasowych we krwi i płynach zewnątrzkomórkowych, dochodzi do bardzo poważnych zakłóceń czynności serca, którego komórki posiadają właściwości elektrofizjologiczne zbliżone do właściwości neuronów. Wzrost stężeń potasu zewnątrzkomórkowego, zgodnie ze wzorem Nernsta, powoduje spadek potencjału spoczynkowego, będącego, jak już wspomniano, potencjałem równowagi dla jonów potasowych. Powstające na tle obniżonego potencjału spoczynkowego potencjały czynnościowe przebiegają nieprawidłowo, co powoduje zaburzenia w przewodzeniu stanu czynnego i w pobudliwości.

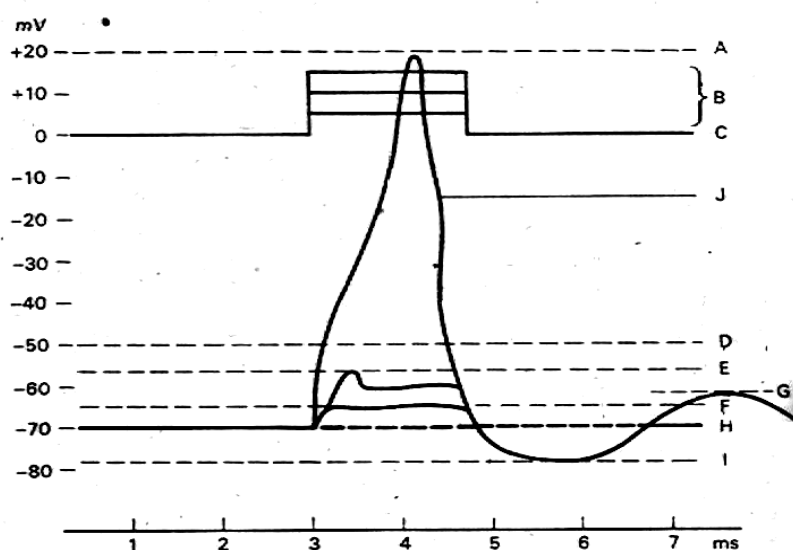
Potencjał czynnościowy i jego geneza

Istnienie potencjału spoczynkowego jest nieodzownym warunkiem dla powstania czynnościowych zjawisk bioelektrycznych w neuronie. Punktem wyjścia dla zjawisk czynnościowych jest zmiana stanu spolaryzowania neuronu. Zmiana ta może iść w kierunku zmniejszenia spoczynkowej różnicy potencjałów, wtedy nosi nazwę depolaryzacji, lub też polegać może na zwiększeniu stanu spolaryzowania i nazywana jest hiperpolaryzacją.

Bardzo pouczające jest prześledzenie poszczególnych etapów powstawania potencjału czynnościowego. Jeżeli na neuron stosuje się elektryczny bodziec depolaryzacyjny, wtedy na oscyloskopowym zapisie wewnątrzkomórkowym zaczną się pojawiać zmiany świadczące o depolaryzacji błony komórki, a więc o zmniejszaniu się potencjału spoczynkowego. Zmiany te, nazywane elektro-tonicznymi, mają charakter bierny i są odzwierciedleniem specyficznej budowy błony neuronu, która nosi w sobie cechy kondensatora. W dalszym ciągu na szczycie zmian elektrotonicznych zaczyna się pojawiać zmiana potencjału, nazywana odpowiedzią

lokalną, która jest już czynną odpowiedzią neuronu na bodziec depolaryzacyjny, choć nie posiada właściwości rozprzestrzeniania się. Gdy opisywane tu zmiany osiągną poziom depolaryzacji krytycznej lub progowej, dochodzi w neuronie do powstania właściwego potencjału czynnościowego, nazwanego ze względu na kształt potencjałem iglicowym (ryc. 24). Jest to potencjał pojawiający się zgodnie z prawem "wszystko albo nic". Oznacza to, że gdy na skutek stymulacji zmiany elektrotoniczne osiągną w neuronie poziom depolaryzacji krytycznej i powstanie potencjał iglicowy, dalsze pobudzanie nie zmienia w danych warunkach amplitudy powstałego potencjału. Istotną cechą potencjału czynnościowego neuronu jest zdolność do rozprzestrzeniania się.

Na ryc. 24 przedstawiono omówione fazy powstawania potencjału czynnościowego oraz jego strukturę. Za główne mechanizmy spoczynkowej różnicy potencjałów uważa się gradienty stężeń jonów po obydwu stronach błony neuronu, przepuszczalność błony dla tych jonów oraz pompę $\text{Na}^+ - \text{K}^+$. Jeżeli przyjmuje się, że potencjał czynnościowy jest zaburzeniem stanu spoczynkowego, musi dojść do zmian wymienionych wyżej parametrów. Neuron nie może zmienić stężeń jonowych z szybkością i w rozmiarach, które by w istotny sposób wpłynęły na wielkość potencjału błonowego. Neuron jest jednak w stanie zmienić właściwości swej błony, a szczególnie przepuszczalność dla jonów. Tak też jest w istocie.



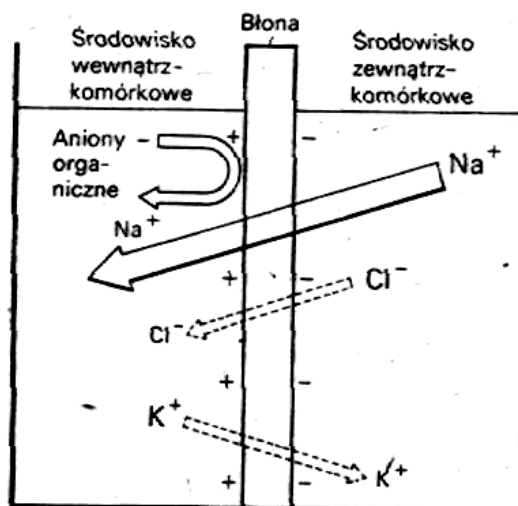
Ryc. 24. Kolejne fazy powstawania potencjału czynnościowego. A — poziom „nadstrzału” (overshoot), B — kolejne impulsy prądu depolaryzacyjnego o wzrastającej amplitudzie, C — poziom zerowy, D — poziom depolaryzacji krytycznej, E — odpowiedź lokalna, F — odpowiedź elektrotoniczna, G — podepolaryzacyjny potencjał następczy, H — poziom potencjału spoczynkowego, I — hiperpolaryzacyjny potencjał następczy, J — potencjał iglicowy.

W myśl teorii zaproponowanej przez A. L. Hodgkina i A. F. Huxleya (1939), potencjał czynnościowy neuronu powstaje w wyniku nagłego, dużego, przejściowego wzrostu przepuszczalności błony neuronów dla Na^+ . Przez bardzo krótki okres przepuszczalność błony neuronu dla Na^+ przewyższa przepuszczalność błony dla innych jonów. Z badań doświadczalnych wynika, że w czasie potencjału czynnościowego przepuszczalność błony neuronu dla Na^+ wzrasta kilkaset razy w porównaniu ze stanem spoczynkowym. Jony Na^+ , wchodząc zgodnie z gradientem stężeń do komórki zmieniają potencjał wnętrza neuronu z ujemnego na dodatni (ryc.25). Podwyższenie przepuszczalności potencjał czynnościowy. Jest to zjawisko refrakcji.

Wkrótce po rozpoczęciu aktywacji Na^+ dochodzi także do wzrostu przepuszczalności błony dla jonów K^+ . Proces ten jest nazywany aktywacją potasową. Szybkość narastania wzrostu przepuszczalności dla jonów K^+ jest znacznie mniejsza niż dla jonów Na^+ . Aktywacja K^+ osiąga swoje maksimum w chwili, gdy przepuszczalność błony dla sodu powraca do wartości wyjściowych (in-aktywacja sodowa). Przepuszczalność błony neuronu dla jonów Na^+ jest największa na szczycie potencjału czynnościowego, po czym ulega szybkiemu zmniejszeniu. W tym obszarze różnica potencjału w poprzek błony neuronu zbliża się do potencjału równowagi dla Na^+ (E_{Na}).

Z tego powodu w wielu neuronach szczyt potencjału iglicowego można przewidzieć na podstawie stężeń sodu na zewnątrz $[Na^+]_z$ z i wewnątrz $[Na^+]_w$ w neuronu w myśl równania Nernsta zastosowanego do sodu

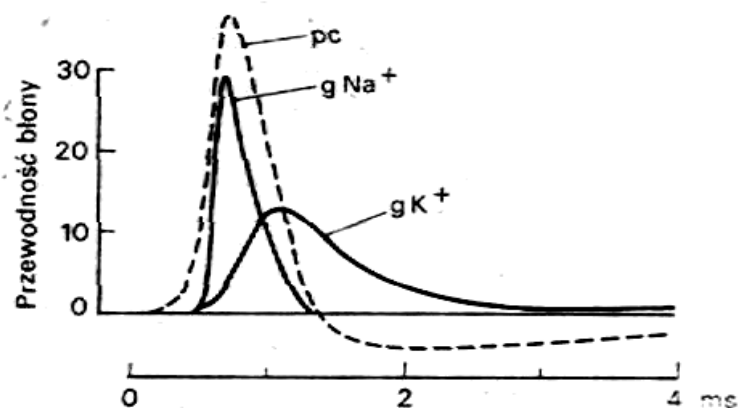
$$E_{Na} = \frac{RT}{F} \ln \frac{[Na^+]_z}{[Na^+]_w} \quad [14]$$



Ryc. 25 Parametry elektrochemiczne w czasie czynności neuronu (zmodyfikowane wg Usherwooda). Przepuszczalność błony dla sodu w czasie czynności neuronu w znaczny sposób przewyższa przepuszczalność dla potasu, a potencjał błonowy zależy głównie od dyfuzji jonów sodowych zgodnie z gradientem stężeń.

Proces inaktywacji sodowej i aktywacji potasowej jest odpowiedzialny za repolaryzację neuronu. Repolaryzacja kończy potencjał czynnościowy. Gdy potencjał równowagi dla potasu nie pokrywa się dokładnie z potencjałem spoczynkowym, wtedy po fazie repolaryzacji przez kilka milisekund potencjał błonowy będzie różnił się od potencjału spoczynkowego. Te przejściowe zjawiska noszą nazwę potencjałów następczych; mogą one mieć charakter zarówno depolaryzacyjny, jak i hiperpolaryzacyjny (ryc. 26). Z przytoczonego opisu genezy iglicowego potencjału czynnościowego wynika, że dla powstania tego potencjału nieodzowna jest obecność jonów Na^+ w środowisku otaczającym neuron.

Wnikliwa analiza opisanych powyżej procesów jonowych, biorących udział w genezie potencjału czynnościowego, stała się możliwa dzięki zastosowaniu nowoczesnej metody stabilizacji potencjału komórkowego (voltage clamp). Metoda ta pozwala na rozdzielenie i zidentyfikowanie prądów jonowych niesionych przez jony Na^+ i K^+ w różnych fazach powstawania potencjału czynnościowego. Okazało się np., że potencjał krytyczny (progowy) błony komórkowej neuronu jest momentem, w którym prąd sodowy jest równy prądowi potasowemu (przy odwrotnych kierunkach).



Ryc. 26 Zmiany przewodności czynnej sodu (g_{Na^+}) i potasu (g_{K^+}) kształtujące potencjał czynnościowy (pc) (zmodyfikowane wg Hodgkina i Huxleya).

Zastosowanie nowoczesnych metod badawczych pozwoliło na opracowanie hipotetycznej koncepcji tzw. kanałów w błonie neuronu, a więc miejsc, w których zlokalizowane są procesy aktywacji i inaktywacji jonowej. Istnieje szereg koncepcji strukturalnych kanałów błonowych. Właściwości kanału związane z przepuszczalnością jonów przez błonę wiążą się z uruchomieniem, podtrzymywaniem czy wreszcie zahamowaniem przechodzenia jonów przez dany kanał. Hipotetyczny kanał błony uważany jest także za mechanizm posiadający zdolność wyboru jonów. W związku z tym istnieją przypuszczenia o istnieniu np. odrębnych kanałów sodowych czy też potasowych.

Punktem przełomowym w badaniach nad mechanizmami aktywacji i inaktywacji jonowej w czasie potencjału czynnościowego było wprowadzenie do tych badań tetrodotoksyny. Jest to jad produkowany przez żyjące głównie w Oceanie Spokojnym ryby z rodziny Te-trodonidae. Jad ten posiada zdolność

wybiórczego blokowania wzrostu przepuszczalności błony neuronu dla jonów Na^+ ; eliminuje to możliwość powstania potencjału czynnościowego. W związku z tym tetrodotoksyna jest jedna z najsilniejszych trucizn. Stosując ją ustalono, że blokowanie wzrostu przepuszczalności błony dla Na^+ polega na łączeniu się jednej cząsteczki tetrodotoksyny z jednym kanałem sodowym. Doświadczenia te pozwoliły na ustalenie liczby miejsc (kanałów) sodowych w błonie neuronu. Liczbę tę szacuje się średnio na około 50 kanałów na $1\mu m^2$ błony.

POBUDLIWOŚĆ, POBUDZENIE I HAMOWANIE

Pobudliwość można określić jako zdolność organizmów, tkanek, czy też komórek do odpowiadania na bodźce. Wykazujące tę zdolność organizmy, tkanki, czy też komórki nazywamy pobudliwymi. Różne organizmy, a w danych organizmach określone tkanki czy też komórki, posiadać mogą różną pobudliwość. Niekiedy te same elementy mogą wykazywać różnice pobudliwości w czasie; w ramach rytmiki biologicznej można obserwować zmiany pobudliwości w ciągu np. doby, miesiąca, czy też roku. W kategoriach całego organizmu odpowiedzi realizowane są jako wynik działania bodźców na specjalnie do tego przystosowane elementy - receptory, będące podstawowym elementem narządów zmysłów. Na poziomie komórki bodźce działają na błonę komórkową.

Określenie poziomu pobudliwości organizmu, tkanek, czy też komórek dokonuje się przede wszystkim metodą pośrednią. Metoda ta polega na pomiarze wielkości stosowanego bodźca. Bardzo pożyteczną, a zarazem powszechnie stosowaną, miarą poziomu pobudliwości jest próg pobudliwości. Przez próg pobudliwości rozumie się najniższy bodziec zdolny do wywołania w danych warunkach określonej reakcji. Jest to bodziec progowy. Bodźce o natężeniu niższym od progowego nazywamy pod-progowymi, o natężeniu wyższym zaś bodźcami nadprogowymi. Ponieważ na organizm, tkankę, czy też komórkę działać mogą bodźce niosące w sobie różne rodzaje energii, bodźce progowe - a tym samym także próg pobudliwości - mogą posiadać różne miana. Można więc wyrażać próg pobudliwości w miliwoitach lub miliamperach (bodziec elektryczny prądów selektywnych), milimolach (bodziec chemiczny) itd.

Jak wynika z przytoczonej tu definicji, pobudliwość przejawia się w zdolność: odpowiadania na bodźce. Komórka nerwowa może odpowiadać na bodźce w rozmaity sposób (np. zmianami częstotliwości potencjałów czynnościowych). Jeżeli jednak jako stan wyjściowy przyjmie się neuron w stanie spoczynku, to na działanie bodźca może on zareagować w dwojaki sposób: zmniejszeniem spoczynkowej różnicy potencjału (depolaryzacją) lub jej zwiększeniem (hiperpolaryzacją). W pierwszym przypadku dochodzi do pobudzenia neuronu, w drugim zaś do jego zahamowania, co można określić z jednej strony jako stan uczynnienia, z drugiej zaś strony jako unieczynnienie neuronu. Zarówno depolaryzacja, jak i hiperpolaryzacja są procesami aktywnymi) wymagającymi energii.

Przewodzenie informacji w obrębie neuronu

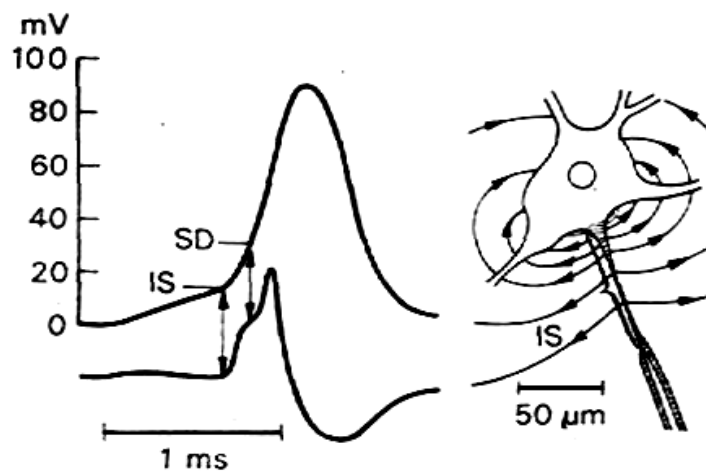
W neuronie bardzo jaskrawo uwidocznione są powiązania i zależności strukturalno-czynnościowe. Wyodrębniono w strukturze neuronu cztery zasadnicze strefy czynnościowe: 1) wejście neuronu (dendryty i częściowo ciało neuronu), 2) strefa inicjacji impulsów (początkowy odcinek aksonu). 3) strefa przewodzenia (głównie akson), 4) wyjście neuronu (zakończenie aksonu). Kolejność, w jakiej podano strefy czynnościowe, jest zarazem obrazem kolejności uczynniania poszczególnych elementów neuronu, a także kierunku przewodzenia informacji w obrębie neuronu.

Dendryty, stanowiące jeden z elementów wejściowych neuronu, charakteryzują się określonymi cechami zarówno strukturalnymi, jak i czynnościowymi, odróżniającymi je od pozostałych części neuronu. W ramach cech strukturalnych na uwagę zasługuje tzw. aparat pączkowy (spine apparatus), znajdujący się w "pączkach" ("kolcach") dendrytów. Aparat składa się z równoległych warstw pęcherzyków. Strukturze tej przypisuje się udział w procesach związanych z torowaniem. Z cech czynnościowych obszaru dendrytycznego wymienić należy znacznie wyższy próg pobudliwości w porównaniu z początkowym odcinkiem aksonu oraz znacznie mniejszą prędkość przewodzenia impulsów w porównaniu z prędkościami występującymi w aksonach.

Przewodzenie informacji od dendrytów do zakończeń aksonu jest kierunkiem fizjologicznym i nosi nazwę przewodzenia ortodromowego. Kierunek przeciwny (nie występujący w organizmach żywych, lecz możliwy do uzyskania eksperymentalnie) nosi nazwę przewodzenia antydromowego.

Na ryc. 27 przedstawiono wewnątrzkomórkowe zapisy zjawisk bioelektrycznych w neuronie ruchowym. Z zapisów tych widać, że ramię wstępujące potencjału iglicowego wykazuje w

swym przebiegu załamki. Jeśli ten charakterystyczny przebieg podda się szczegółowej analizie poprzez zastosowanie różniczkowania elektrycznego, wtedy wyraźnie widać, że na właściwy potencjał czynnościowy składają się oddzielone od siebie dwa elementy bioelektryczne. Te dwa elementy są obecnie określane ze względu na miejsce powstawania w neuronie jako iglica odcinka początkowego aksonu wraz ze wzgórkim aksonu oraz iglica ciała komórki i dendrytów. Badania elektrofizjologiczne wykazały istnienie różnic w poziomie depolaryzacji krytycznej w rejonach odcinka początkowego aksonu oraz ciała komórki i dendrytów. Jeśli chodzi o dotychczas zbadane neurony ruchowe, to próg odcinka początkowego aksonu jest zawsze niższy od progu ciała komórki i dendrytów i co więcej, stosunek ten wynosi co najmniej 1 : 2 tzn., że próg ciała komórki i dendrytów jest co najmniej dwukrotnie wyższy od progu odcinka początkowego aksonu. Stymulacja wejścia neuronu (dendryty i ciało neuronu) nie prowadzi bezpośrednio do powstania potencjału czynnościowego w tym rejonie, a jedynie do szerzenia się fali depolaryzacji w kierunku odcinka początkowego aksonu, który ze względu na niski próg jest miejscem wyzwolenia potencjału czynnościowego (ryc.27).

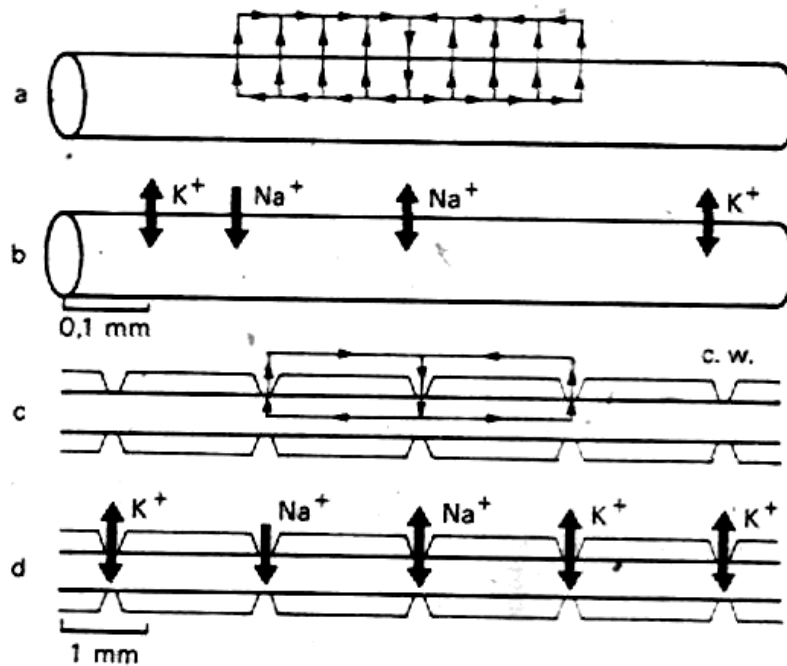


Ryc. 27 Wewnątrzkomórkowe zapisy potencjałów czynnościowych neuronu ruchowego. IS — iglica początkowego odcinka aksonu, SD — iglica ciała komórki i dendrytów (zmodyfikowane wg Coombsa i wsp.).

Główną strefą przewodzenia fali depolaryzacji w neuronie jest akson. W przewodzeniu aksonalnym zaznaczają się wyraźnie ściśle związki między strukturą a funkcją. Po wyzwoleniu we włóknie nerwowym potencjału czynnościowego w miejscu wyzwolenia nie może w okresie ok. 1 ms powstać następna odpowiedź. Okres, w którym nie może powstać - bez względu na natężenie stosowanego bodźca - kolejny potencjał czynnościowy, nazywany jest okresem refrakcji bezwzględnej.

W przebiegu powstawania potencjału czynnościowego okres refrakcji bezwzględnej odpowiada okresowi inaktywacji sodowej, a więc -stopniowego zmniejszania przepuszczalności błony dla jonów Na^+ . Po okresie refrakcji bezwzględnej następuje okres, w którym neuron odzyskuje wprawdzie swą pobudliwość, lecz próg depolaryzacji krytycznej jest podwyższony (do wywołania potencjału czynnościowego konieczny jest silniejszy bodziec). Ten okres trwający do kilku milisekund (w zależności od rodzaju włókien) nazywany jest okresem refrakcji względnej. Zjawisko refrakcji przeciwdziała nakładaniu się potencjałów czynnościowych na siebie; w związku z tym występują one jako odrębne zjawiska bioelektryczne. Należy także zdać sobie sprawę z tego, że okresy refrakcji determinują w dużej mierze częstotliwość wyładowań neuronu.

Sposób przewodzenia informacji we włóknach nerwowych zależy głównie od tego, czy włókna te mają osłonkę mielinową, czy też są jej pozbawione. We włóknach bez osłonki mielinowej potencjały czynnościowe wędrują ruchem jednostajnym ze stałą dla danych warunków prędkością. Ten typ przewodzenia nosi nazwę przewodzenia ciągłego. We włóknach z osłonką mielinową



Ryc. 28 Schemat przewodzenia potencjałów czynnościowych wzdłuż aksonu (zmodyfikowane wg Usherwooda). Przewodzenie nerwowe w aksonie bez osłonki mielinowej (a, b) oraz w aksonie z osłonką mielinową (c, d). Prądy elektryczne związane z potencjałami czynnościowymi typu „wszystko albo nic” przedstawiono na ryc. a i c, podczas gdy ruchy jonów przedstawiono na ryc. b i d. Strzałki jednostronne wskazują ruch jonów netto, a strzałki dwustronne przedstawiają jon w stanie równowagi elektrochemicznej. C.W. — cieśń węzła.

przewodzenie potencjałów czynnościowych odbywa się z niejednorodną prędkością. Prędkość przewodzenia jest bardzo duża w obszarach pomiędzy cieśniami węzłów, natomiast w samych cieśniach węzłów następuje "przestój" potencjału czynnościowego. Ten rodzaj przewodzenia nazywa się przewodzeniem skokowym. Zasady działania mechanizmów przewodzenia ciągłego i skokowego obrazuje ryc. 28.

Ze względu na stałe procesy odnowy jonowej (w przewodzeniu ciągłym wzdłuż całego włókna, a w przewodzeniu skokowym w cieśniach węzłów) amplituda potencjału czynnościowego nie maleje w miarę przesuwania się tego potencjału wzdłuż włókien nerwowych. Jest to zjawisko zwane przewodzeniem bez dekrementu (strat). W niektórych przypadkach, szczególnie przy zmianach chorobowych, występuje przewodzenie z dekrementem, co oczywiście zaburza czynność układu nerwowego.

Prędkość przewodzenia impulsów wzdłuż włókien nerwowych zależy od średnicy włókien nerwowych oraz od sposobu przewodzenia. Prędkość przewodzenia potencjału czynnościowego zależy od podłużnej oporności włókna nerwowego. Ta oporność, podobnie jak w układach nieożywionych, jest odwrotnie proporcjonalna do kwadratu średnicy przewodnika. Wynika stąd, że włókna o dużym przekroju będą miały względnie niską oporność aksoplazmy i co za tym idzie - zdolność do większej prędkości przewodzenia potencjałów czynnościowych, włókna zaś o przekroju mniejszym będą miały wyższą oporność aksoplazmy i przewodzić będą wolniej. Obok wyraźnego wpływu wielkości przekroju włókna na prędkość

przewodzenia fali depolaryzacji, ważnym czynnikiem ją kształtującym jest sposób przewodzenia potencjałów. We włóknach rdzennych, dzięki dobrym właściwościom izolacyjnym osłonki mielinowej, prędkość przewodzenia jest znacznie wyższa, niż wskazywałby na to porównawczo przekrój włókna. Prędkość przewodzenia fali depolaryzacji we włóknach nerwowych jest "dopasowana" do funkcji danych włókien.

Grupy włókien nerwowych

Grupa	A				B	C	
Podgrupa	α	β	γ	δ	-	s	d.r.
Średnica aksonu w μm	12-20	5-12	3-6	2-5	+/-3	0,3-1,3	0,4-1,2
Prędkość przewodzenia w m/s	70-120	30-70	15-30	12-30	3-15	0,7-2,3	0,5-2,0
Oślonka mielinowa	+	+	+	+	+	-	-
Włókna nerwowe aferentne	+	+		+			+
Włókna nerwowe eferentne	+		+		+	+	
Czas trwania potencjału czynnościowego w ms	0,4-0,5				1,2	2	2
Okres refrakcji bezwzględnej w ms	0,4-1,0				1,2	2	2

a więc do biologicznego znaczenia szybkości reakcji.

Na podstawie cech zarówno morfologicznych, jak i czynnościowych dzieli się włókna nerwowe na cztery następujące grupy:

1) Włókna nerwowe grupy A, do których należą włókna z osłonką mielinową, spełniające rolę zarówno aferentną (dośrodkową), jak i eferentną (odśrodkową). W grupie tej - biorąc pod uwagę średnicę włókien - wyróżnia się cztery podgrupy (α, β, γ i δ).

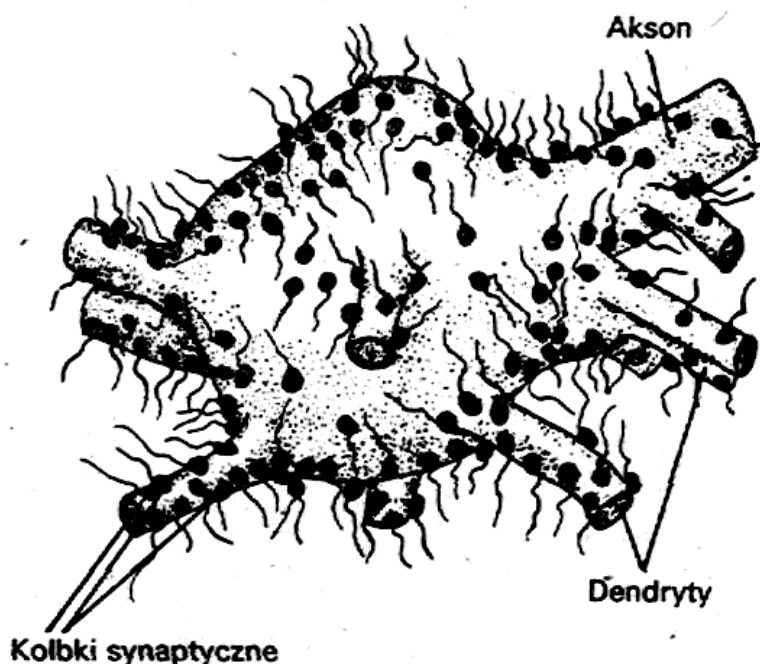
2) Włókna grupy B są częścią układu autonomicznego. Posiadają one osłonkę mielinową i są włóknami przywspółczulnymi oraz współczulnymi przedzwojowymi.

3) Włókna grupy Cs (sympathetic C fibers) - to włókna współczulne zazwojowe, pozbawione osłonki mielinowej.

4) Włókna grupy C d.r. (dorsal root C fibers) - to włókna wstępujące do rdzenia przez korzenie grzbietowe. Są one pozbawione osłonki mielinowej. Szczegółowe dane dotyczące powyższych grup włókien przedstawiono w tabeli

Przekazywanie informacji innym komórkom

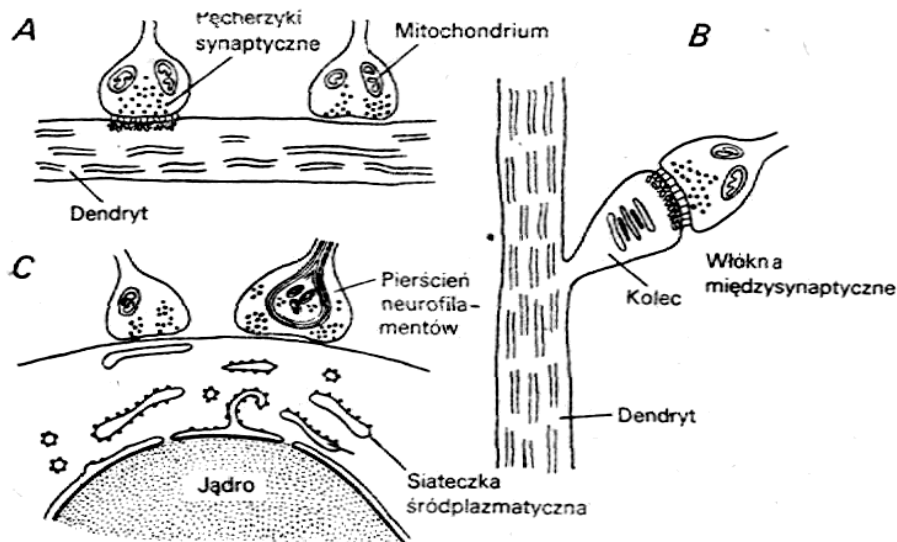
Informacja nerwowa, która dociera do wyjścia neuronu, a więc do zakończeń aksonów, musi być, zgodnie z założeniami czynnościowymi układu nerwowego, przekazana innym komórkom nerwowym, lub też efektorom, (np. mięśnie, gruczoły). Istnienie tej ciągłości czynnościowej jest nieodzownym warunkiem prawidłowego funkcjonowania układu nerwowego, a tym samym, oczywiście, całego organizmu. W parze ze wspomnianą ciągłością czynnościową nie idzie ciągłość strukturalna. Poszczególne komórki nerwowe nie są ze sobą zespolone, a jedynie stykają się ze sobą. To samo można powiedzieć o sposobie łączenia się zakończeń aksonu z efektorom. Ten obszar styku nosi nazwę synapsy. W zależności od rodzaju stykających się elementów wyróżnia się synapsy nerwowo-nerwowe (styk dwóch neuronów), synapsy nerwowo-mięśniowe. Połączenia między dwiema komórkami nerwowymi mogą przebiegać rozmaicie i w związku z tym wyróżnia się szereg synaps, których nazwy wskazują na charakter styku międzyneuralnego. Wyróżnia się więc synapsy aksodendrytyczne, aksoaksonalne i aksosomatyczne (ryc. 29).



Ryc. 29 Model komórki przedniego rogu rdzenia kręgowego. Czarne plamki oznaczają kolbki synaptyczne neuronów presynaptycznych (wg Ganonga).

W skład synapsy wchodzi trzy zasadnicze elementy strukturalne: element presynaptyczny, szczelina synaptyczna i element postsynaptyczny. Element presynaptyczny tworzą zakończenia aksonów. Błona zakończeń jest nazywana błoną presynaptyczną. Szczelina synaptyczna jest przerwą strukturalną między elementami pre- i postsynaptycznymi. W skład elementu postsynaptycznego w synapsie między dwoma neuronami mogą wchodzić różne części neuronu postsynaptycznego (np. dendryty, ciało neuronu), w zależności od charakteru synapsy. W związku z lokalizacją kluczowych procesów postsynaptycznych w błonie elementu postsynaptycznego wprowadzono pojęcie błony postsynaptycznej.

Najbardziej istotnym momentem w przekazywaniu informacji innym komórkom poprzez synapsy jest zmiana nośnika dla informacji. W elemencie presynaptycznym nośnikiem dla przesyłania informacji są potencjały czynnościowe. W obrębie natomiast samej synapsy dochodzi do zmiany nośnika elektrycznego na chemiczny. Dzieje się to na skutek wydzielania przez element presynaptyczny substancji chemicznych zwanych mediatorami synaptycznymi (przebieżniki synaptyczne, substancje przebieżnikowe). W obrębie synapsy właśnie te substancje przekazują nadaną przez element presynaptyczny informację nerwową do elementu postsynaptycznego. Synapsy posługujące się takim właśnie sposobem przekazywania informacji nazywane są synapsami chemicznymi.



Ryc. 30 Różne rodzaje połączeń synaptycznych (wg Ganonga). A — dwa rodzaje połączeń z dendrytami (synapsy aksodendrytyczne), B — synapsa na kolcu dendrytu. C — połączenia aksosomatyczne.

Rola synaps i mediatorów chemicznych

Dochodzący - zgodnie z kierunkiem ortodromowym - do zakończeń aksonu potencjał czynnościowy jest mechanizmem spustowym dla zapoczątkowania procesu przenoszenia synaptycznego. Pierwszym etapem tego procesu jest tzw. sprzężenie elektrowydzielnicze. Jest to proces polegający na zainicjowaniu przez falę depolaryzacji uwalniania mediatora synaptycznego do szczeliny synaptycznej. Dla zabezpieczenia ciągłości czynnościowej w obrębie układu nerwowego konieczne jest przejście mediatora poprzez szczelinę synaptyczną i połączenie z elementem postsynaptycznym. To połączenie odbywa się na poziomie błony postsynaptycznej, a ściślej, na poziomie receptora błonowego, specyficznego dla danego mediatora synaptycznego. W błonie postsynaptycznej ma miejsce sprzężenie chemiczno-elektryczne. Oznacza to, że informacja znów korzystać będzie z nośnika elektrycznego. W następstwie połączenia się mediatora synaptycznego z receptorem¹ błonowym powstają zmiany przepuszczalności błony postsynaptycznej dla jonów. Wynikiem tych zmian jest powstanie zjawisk bioelektrycznych, które ogólnie nazywamy potencjałami postsynaptycznymi. Potencjały postsynaptyczne należą do grupy potencjałów nie rozprzestrzeniających się, a także stopniowanych. Ta ostatnia cecha, która jest przeciwieństwem mechanizmu typu "wszystko albo nic", oznacza, że amplituda potencjału zależy od natężenia bodźca. Potencjały postsynaptyczne mogą mieć charakter depolaryzacyjny, lub też hiperpolaryzacyjny. W pierwszym przypadku mogą one doprowadzić do depolaryzacji krytycznej i w rezultacie do wyzwolenia potencjału czynnościowego i dlatego mają charakter pobudzający. W związku z tym depolaryzacyjne potencjały postsynaptyczne nazywa się postsynaptycznymi potencjałami pobudzającym: (EPSP - excitatory postsynaptic potentials). Postsynaptyczne potencjały hiperpolaryzacyjne poprzez zwiększenie różnicy potencjałów wpływają w sposób hamujący na błonę postsynaptyczną. Nazywa się je dlatego postsynaptycznymi potencjałami hamującymi (IPSP - inhibitory postsynaptic potentials). Te dwa rodzaje wyładowań postsynaptycznych decydują o przynależności danej synapsy do kategorii pobudzających, czy też hamujących.

W związku z procesami synaptycznymi wpływa problem pobudliwości błony neuronu w odniesieniu do specyficznych bodźców. Cała nieomalże powierzchnia błony neuronu wykazuje pobudliwość elektryczną. Ta pobudliwość nie jest jednakowa we wszystkich częściach neuronu. Najwyższą pobudliwość elektryczną (najniższy próg pobudliwości) wykazuje odcinek początkowy aksonu. Istnieją wszakże obszary błony neuronu, które należy uznać za niepobudliwe elektrycznie. Są to te obszary błony neuronu postsynaptycznego, które wchodzi w skład synapsy chemicznej. Nie odpowiadają one na bodźce elektryczne, natomiast wykazują pobudliwość chemiczną (zdolność do reakcji na działanie mediatora synaptycznego).

Innym parametrem związanym z czynnością synaps jest tzw. opóźnienie synaptyczne. Jest to czas, jaki jest potrzebny na przejście informacji nerwowej poprzez synapsę. W zależności od rodzaju synapsy i warunków, opóźnienie synaptyczne może być zawarte w granicach od 0,5 do kilku milisekund. Sama nazwa „opóźnienie” sugeruje, że synapsa opóźnia, zwalnia przewodzenie informacji. Tak jest w istocie. Prędkość przechodzenia informacji poprzez synapsę jest niewspółmiernie mniejsza od prędkości przewodzenia wzdłuż włókien nerwowych. Jest to zrozumiałe ze względu na odmienność podłoża i charakter przewodzenia.

Połączenia synaptyczne pomiędzy neuronami są obszarem aktywnej działalności integracyjnej. Dzieje się tak dlatego, że potencjały synaptyczne są zjawiskami elektrycznymi stopniowanymi, w odróżnieniu od zjawisk typu „wszystko albo nic” występujących w aksonach. Z tego powodu istnieje możliwość oddziaływania na siebie i łączenia zarówno w czasie, jak i przestrzeni. Nawet pojedyncza synapsa może być obszarem integracji poprzez sumowanie i torowanie potencjałów synaptycznych. Jeżeli weźmie się jeszcze pod uwagę, że większość neuronów w organizmie odbiera jednocześnie impulsy z synaps pobudzających i hamujących, wtedy w pełni zdać można sobie sprawę z możliwości integracji informacji nerwowej w obrębie synaps.

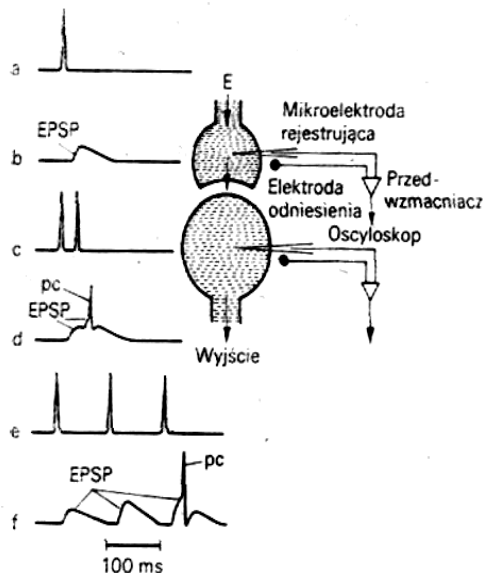
Do podstawowych procesów, będących wyrazem integracji na poziomie synaps, należą: sumowanie (w czasie i przestrzeni), torowanie, hamowanie presynaptyczne, hamowanie postsynaptyczne.

W niektórych synapsach każdy pre-synaptyczny potencjał czynnościowy wywołuje postsynaptyczny potencjał czynnościowy. W takim wypadku mówi się o przenoszeniu 1 : 1. Jeżeli jednak potencjały postsynaptyczne są zjawiskami podprogowymi (nie osiągającymi poziomu depolaryzacji krytycznej dla wyzwolenia potencjału czynnościowego), wtedy postsynaptyczny potencjał czynnościowy powstanie po sumowaniu w czasie lub po torowaniu pobudzających potencjałów postsynaptycznych lub też po obydwu tych procesach. W sumowaniu postsynaptyczne potencjały pobudzające dodają się, natomiast podczas torowania wielkość potencjałów synaptycznych rośnie wraz z powtórzeniem bodźca. Tak może działać w bardzo prostym przypadku pojedynczej synapsy. Kiedy neuron postsynaptyczny otrzymuje informację od większej liczby komórek, wtedy możliwości integracji stają się większe. Obok sumowania w czasie występować może sumowanie przestrzenne.

Procesy integracyjne w obszarze synaptycznym stają się jeszcze bardziej złożone, gdy neuron posiada wejście zarówno pobudzające, jak i hamujące. Wtedy dochodzi do sumowania potencjałów postsynaptycznych zarówno pobudzających, jak i hamujących. Jest to hamowanie postsynaptyczne. Integracja synaptyczna komplikuje się jeszcze bardziej przy zjawisku hamowania pre-synaptycznego. Ten rodzaj hamowania związany jest z istnieniem synaps akso-aksonalnych leżących w rejonie presynaptycznym akso-dendrytycznej lub akso-somatycznej synapsy pobudzającej. Intensywne badania mikroelektrodowe doprowadziły do ustalenia dwóch różnych mechanizmów hamowania presynaptycznego. W pierwszym z nich elementem presynaptycznym synapsy akso-aksonalnej jest neuron hamujący, którego działalność doprowadza do hiperpolaryzacji części presynaptycznej synapsy akso-dendrytycznej lub akso-somatycznej. Hiperpolaryzacja ta uniemożliwia przewodzenie impulsów w tym rejonie, w wyniku czego nie dochodzi do uwolnienia mediatora, a co za tym idzie, przenoszenie zostaje zablokowane.

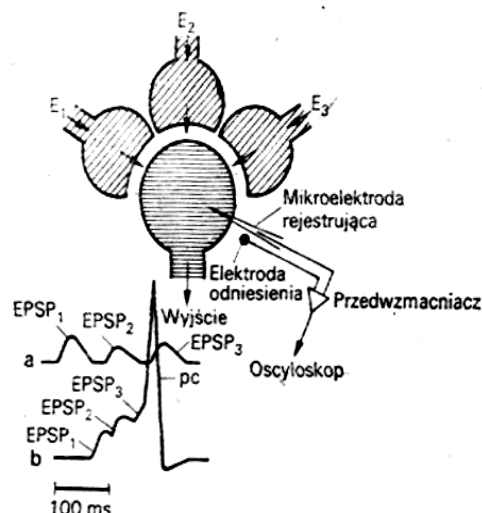
Inny rodzaj hamowania presynaptycznego może być realizowany na podłożu synapsy akso-aksonalnej o charakterze pobudzającym. W tym przypadku działalność tej synapsy doprowadza do depolaryzacji jej postsynaptycznej części. Depolaryzacja powoduje zmniejszenie amplitudy potencjałów czynnościowych biegnących w aksonie synapsy akso-dendrytycznej lub akso-somatycznej. Ilość uwalnianego mediatora zależy nie tylko od częstotliwości, ale także od amplitudy potencjałów czynnościowych dochodzących do rejonu pre-synaptycznego. Opisane zmniejszenie amplitudy potencjałów czynnościowych doprowadza więc w rezultacie do zmniejszenia ilości uwalnianego mediatora, co może spowodować zahamowanie procesu przenoszenia synaptycznego. Ten typ hamowania presynaptycznego charakteryzuje się działaniem to-nicznym, tzn. może utrzymywać się przez dłuższy okres (do 200 ms).

Opisane tu mechanizmy integracji na poziomie synaps przedstawiono i objaśniono bardziej szczegółowo na ryc. 31-34.



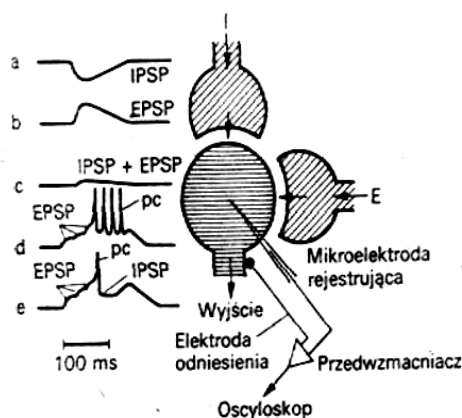
Ryc. 31

Ryc. 31 Sumowanie w czasie oraz torowanie (zmodyfikowane wg Usherwooda). a, c, e — potencjały czynnościowe neuronu presynaptycznego (E), b, d, f — odpowiedzi neuronu postsynaptycznego, b — EPSP nie osiąga poziomu depolaryzacji krytycznej, d — sumowanie się EPSP prowadzi do wyzwolenia potencjału czynnościowego (pc) (sumowanie w czasie), f — na skutek wzrostu amplitudy EPSP przy powtarzaniu stymulacji następuje wyzwolenie potencjału czynnościowego (pc) (torowanie).



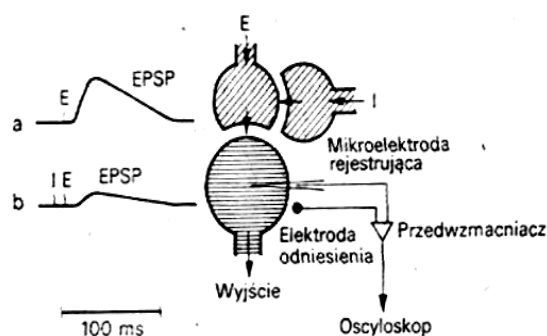
Ryc. 32

Ryc. 32 Sumowanie przestrzenne (zmodyfikowane wg Usherwooda). Stymulacja każdego z pobudzających neuronów presynaptycznych (E_1 , E_2 , E_3) wytwarza potencjał postsynaptyczny, który jest niedostateczny do wywołania potencjału czynnościowego w komórce postsynaptycznej (a). Jednakże, jeśli wszystkie trzy wejścia synaptyczne podlegają stymulacji równocześnie, wtedy potencjały postsynaptyczne sumują się i depolaryzacja komórki postsynaptycznej jest dostateczna do wyzwolenia potencjału czynnościowego (b).



Ryc. 33

Ryc. 33 Hamowanie postsynaptyczne (zmodyfikowane wg Usherwooda). Stymulacja wejścia pobudzającego (E) prowadzi do powstania podprogowego EPSP (b), stymulacja zaś wejścia hamującego (I) do powstania IPSP (a). Równoczesna stymulacja E oraz I prowadzi do algebraicznego sumowania się potencjałów postsynaptycznych (c). Pokazano (e) hamujący wpływ IPSP na salwę potencjałów czynnościowych (pc), wyzwolonych jako wynik sumowania (d).



Ryc. 34

Ryc. 34 Hamowanie presynaptyczne przy udziale neuronu hamującego (zmodyfikowane wg Usherwooda). Stymulacja wejścia pobudzającego (E) prowadzi do powstania w neuronie postsynaptycznym EPSP (a). Równoczesna stymulacja wejścia pobudzającego (E) i hamującego (I) prowadzi do przytłumienia EPSP (b).

Biosynteza, transport wewnątrzkomórkowy, magazynowanie i uwalnianie mediatorów synaptycznych pobudzających i hamujących

W 1958 roku Paton zaproponował następujące kryteria, które winna spełniać substancja chemiczna, aby mogła być zaklasyfikowana jako mediator chemiczny:

1) Substancja ta winna być zawarta w neuronie presynaptycznym, który winien posiadać zdolność do jej syntezy.

2) Na skutek stymulacji aksonu presynaptycznego substancja ta powinna być uwalniana.

3) Podziałanie tą substancją na komórkę postsynaptyczną powinno prowadzić do efektu charakterystycznego dla normalnego przenoszenia poprzez badaną synapsę.

4) Działanie badanej substancji na komórkę postsynaptyczną powinno podlegać wpływom środków blokujących w ten sam sposób, w jaki ich wpływom podlega normalne przenoszenie synaptyczne.

5) W okolicy działania mediatora powinien znajdować się enzym zdolny do jego unieczynnienia przez rozkład.

Piąte kryterium ukazuje jedynie jeden ze sposobów unieczynnienia mediatora chemicznego. Obok rozkładu enzymatycznego mediator synaptyczny może być usuwany ze szczeliny synaptycznej przez dyfuzję lub reabsorpcję.

Cały proces działania mediatorów w obrębie synapsy można podzielić na kolejne, zalegające się etapy: biosynteza, transport wewnątrzkomórkowy, magazynowanie, uwalnianie, łączenie się z receptorem błony postsynaptycznej, unieczynnianie. Synteza mediatorów synaptycznych zachodzić może w zasadzie w całym obszarze neuronu, choć zintensyfikowanie tego procesu obserwuje się przede wszystkim w zakończeniach aksonu. Obecnie nauka dysponuje już danymi dotyczącymi szczegółowych cykli syntezy mediatorów chemicznych, w których to cyklach wiodącą rolę odgrywają układy enzymatyczne. W zależności od miejsca syntezy wpływa sprawa wewnątrzkomórkowego transportu mediatora. W 1968 roku opublikowana została przez Schmitta teoria o transporcie tzw. pęcherzyków synaptycznych (zawierających mediator) od ciała neuronu wzdłuż aksonu do jego zakończeń. W procesie tym ważną rolę odgrywają neurotubule aksonalne, wzdłuż których odbywa się transport wewnątrzkomórkowy.

Miejscem magazynowania przekaźnika synaptycznego są przede wszystkim pęcherzyki synaptyczne. Pęcherzyki te o średnicy 20-60 nm są częścią składową 'kolbek synaptycznych'. Te ostatnie są strukturami tworzącymi zakończenie aksonu.

Proces uwalniania mediatora synaptycznego ma kilka charakterystycznych cech. Chemiczne przenoszenie synaptyczne może dojść do skutku jedynie wtedy, gdy istnieje substancja przekaźnikowa gotowa do uwolnienia do szczeliny synaptycznej w chwili, gdy błona presynaptyczna zostanie zdepolaryzowana. Pierwszą z charakterystycznych cech uwalniania mediatora synaptycznego jest uwalnianie tzw. pakietów mediatora, a więc pewnej, określonej liczby cząsteczek przekaźnika (quantal release). Fakt ten jest poparciem dla hipotezy pęcherzykowej, gdyż wtedy pakiet mediatora można sobie wyobrazić jako zawartość jednego pęcherzyka synaptycznego. Zakłada się np., że w jednym pęcherzyku synaptycznym synapsy cholinergiczej znajduje się ok. 1000 cząsteczek acetylocholiny.

Drugą cechą charakterystyczną uwalniania mediatora jest to, że mediator uwalnia się z zakończeń aksonu przez cały czas, niezależnie od tego, czy do zakończeń aksonu dociera impuls - czy nie. Fakt ten może być w pełni zrozumiały, jeśli się doda, że między uwalnianiem "spoczynkowym" a "czynnościowym" istnieją duże różnice ilościowe. W czasie, gdy do zakończeń aksonu nie dociera impulsacja nerwowa, z zakończeń tych uwalniają się spontanicznie małe ilości mediatora. Istnieje przypuszczenie, że zjawisko to można przypisać faktowi, iż pęcherzyki synaptyczne znajdując się w ciągłym ruchu, w pewnych momentach stykają się z "obszarem uwalniania" (tzw. obszar aktywny) i uwalniają swą zawartość do szczeliny synaptycznej. Depolaryzacja błony presynaptycznej przez potencjał czynnościowy prowadzi do bardzo znacznego zwiększenia „obszaru uwalniania”, wzrasta wówczas możliwość zetknięcia się pęcherzyków synaptycznych z takim obszarem. W ten sposób do szczeliny synaptycznej uwolniona zostaje równocześnie duża liczba "pakietów".

Trzecią wreszcie charakterystyczną cechą procesu uwalniania mediatora jest rola jonów wapniowych w tym procesie. Okazało się, że pod nieobecność jonów Ca^{++} w płynie

zewnątrzkomórkowym depolaryzacja aksonu przez potencjał czynnościowy nie doprowadza do zwiększonego uwalniania mediatora. W takich warunkach spontaniczne uwalnianie małych ilości mediatora nie zostaje zatrzymane. Można z tego wysunąć wniosek, że Ca^{++} jest niezbędny w procesie uruchamiania dodatkowych "obszarów uwalniania" w bio-nie presynaptycznej. Mechanizmem spustowym aktywacji wapnia (wzrost przepuszczalności błony presynaptycznej dla tego jonu) jest depolaryzacja obszaru presynaptycznego. Podobne efekty do spowodowanych niedoborem lub nieobecnością Ca^{++} obserwuje się, stosując na obszar synaptyczny zwiększone stężenie Mg^{++} . Hamujące działanie Mg^{++} na przenoszenie synaptyczne tłumaczone jest kompetytywnym działaniem tych jonów na obszary błony presynaptycznej "zarezerwowane" dla jonów wapniowych.

Wydzielony z części presynaptycznej mediator dyfunduje poprzez szczelinę synaptyczną w kierunku błony postsynaptycznej. Następuje następny etap działalności mediatora: łączenie się z receptorem błony postsynaptycznej. Receptory błonowe są zdefiniowanymi chemicznie obszarami dużych cząsteczek, które łączą się z mediatorem na zasadzie "chemicznego dopełnienia". Jako wynik połączenia się mediatora z receptorem błonowym powstają zmiany przepuszczalności błony postsynaptycznej dla różnych jonów, a w dalszej kolejności zjawiska bioelektryczne, nazywane ogólnie zjawiskami postsynaptycznymi.

Spontaniczne uwalnianie mediatora powoduje powstanie w błonie postsynaptycznej tzw. miniaturowych potencjałów postsynaptycznych. Potencjały te, o amplitudzie ok. 0,7 mV, stwierdzono i opisano najpierw na preparatach nerwowo-mięśniowych. Okazało się jednak, że i synapsy nerwowo-nerwowe są miejscem powstawania miniaturowych potencjałów postsynaptycznych. Potencjały te uważa się za wynik uwalniania pojedynczych pakietów mediatora.

Zwiększone uwalnianie mediatora na skutek depolaryzacji błony presynaptycznej prowadzi do powstania w części postsynaptycznej typowych postsynaptycznych potencjałów pobudzających (EPSP), lub też postsynaptycznych potencjałów hamujących (IPSP), w zależności od charakteru synapsy. Pobudzające potencjały postsynaptyczne dają początek potencjałom czynnościowym.

W 1970 r. dwaj badacze brytyjscy Katz i Miledi opisali jeszcze jedno postsynaptyczne zjawisko bioelektryczne. Pracując na preparacie nerwowo-mięśniowym żaby, wyodrębnili oni w błonie postsynaptycznej potencjały o amplitudzie 0,3 V, które według ich hipotezy mają być odzwierciedleniem połączenia się jednej cząsteczki mediatora z receptorem cholinergicznym. Potencjały te nazwane zostały szumem acetylocholinowym.

Ostatnim wreszcie etapem losów mediatora w obrębie synapsy jest jego unieczynnienie. Jednym ze sposobów unieczynnienia mediatora jest jego rozkład za pomocą enzymu. Taki los spotyka np. acetylcholinę, która jest rozkładana przez acetylcholinesterazę na cholinę i kwas octowy. Istnieją wszakże inne, bardziej powszechne, sposoby unieczynniania mediatorów synaptycznych. Do nich należy wychwyt mediatora przez elementy presynaptyczne, czy też gładź, jak to ma miejsce np. w przypadku amin katecholowych, kwasu gamma-aminomasłowego, czy też glicyny.

Mimo intensywnych badań w zakresie identyfikacji mediatorów synaptycznych, w wielu obszarach układu nerwowego nie udało się ustalić, jakie substancje pełnią rolę przenoszącą. Spośród mediatorów pobudzających najbardziej zasadniczą rolę spełnia acetylcholina (ACh). Istnieje cały szereg związków, którym poza acetylcholiną przypisywane są właściwości chemicznych mediatorów stanów pobudzających. Spośród mediatorów hamujących, wywołujących hiperpolaryzację neuronu postsynaptycznego, na pierwsze miejsce wysuwa się kwas gamma-aminomasłowy (GABA). Znaczna część synaps w ośrodkowym układzie nerwowym gromadzi ten związek.

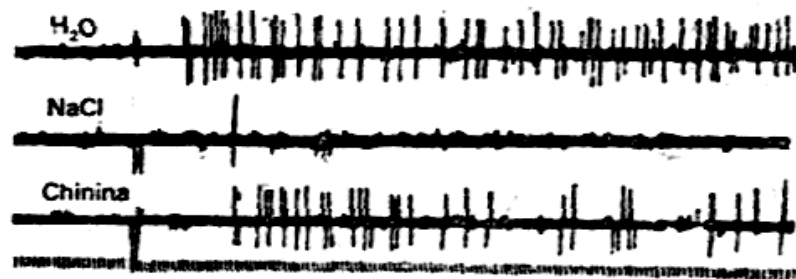
W okresie kilkudziesięciu lat ogólnie przyjęty był pogląd, że jeden neuron może syntetyzować i uwalniać ze swych zakończeń jedynie jeden mediator chemiczny. Jest to tzw. prawo Dale'a, które wraz z Loewim w 1936 r. otrzymał nagrodę Nobla za prace z zakresu chemicznych mediatorów synaptycznych.

Najnowsze badania wykazują jednak, że występują neurony syntetyzujące dwa mediatory synaptyczne.

KODOWANIE INFORMACJI PRZEZ KOMÓRKI NERWOWE

Bodźce działające na układ nerwowy zawierają informacje pochodzące zarówno z wnętrza organizmu, jak i z jego otoczenia. Bodźce te reprezentują rozmaite formy energii. Prawidłowa działalność układu nerwowego, a co za tym idzie, całego organizmu, wymaga w związku z poruszonym tu problemem obecności dwóch procesów. Pierwszym jest proces przetworzenia energii związanej z danym bodźcem na impulsację elektryczną, a więc na zjawiska bioelektryczne. Drugim zaś procesem jest kodowanie odebranej informacji w obrębie układu nerwowego.

Wydawałoby się, że jednym z najprostszych sposobów kodowania informacji w obrębie układu nerwowego mógłby być kod amplitudy. Układ nerwowy nie korzysta jednak z tego sposobu informacji (por. prawo "wszystko albo nic"). Informacja biegnąca w układzie nerwowym zakodowana jest za pomocą kodu częstości. W tym rodzaju kodowania miarą przenoszonego parametru informacji jest gęstość impulsów. W odróżnieniu od kodu interwałowego w kodzie częstości nie chodzi o przerwy między impulsami, lecz o średnią częstość impulsów w określonym przedziale czasu. Jak widać z założeń tego sposobu kodowania, zmiana amplitudy potencjałów, lub też "zgubienie" pojedynczych impulsów nie może wpłynąć w istotny sposób na treść zakodowanej informacji. Według współczesnych poglądów, pojemność informacyjna (zdolność przenoszenia) danego kanału komunikacyjnego rozpatrywana jest w kategoriach szybkości przekazywania informacji. Według definicji pojemność informacyjna kanału to maksymalna ilość informacji przesyłanych w jednostce czasu (bit/s).



Ryc. 35 Reakcje pojedynczego włókna struny bębenkowej na różne roztwory opłukujące język (wg Gawrońskiego, 1966).

Spośród koncepcji dotyczących określania pojemności informacji w kanałach nerwowych na uwagę zasługuje następująca, wyrażona wzorem:

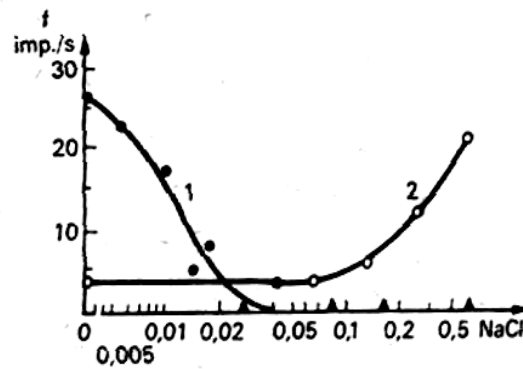
$$C = \frac{1}{T_i} \lg T_i \int \max + 1 \quad [15]$$

gdzie:

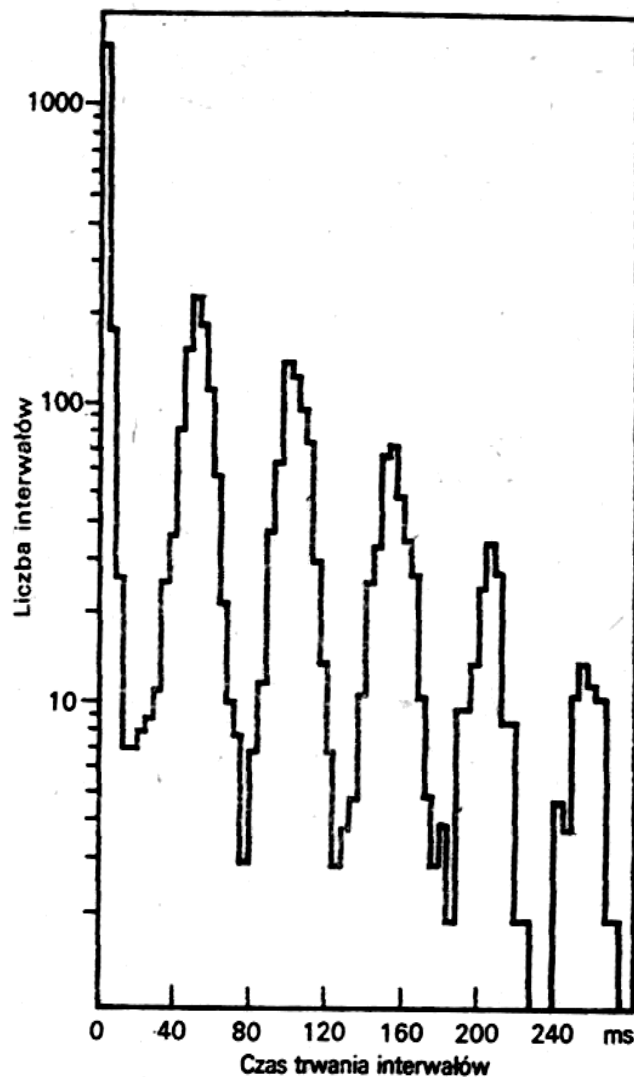
C — średnia pojemność informacyjna,
 $T_i \int \max$ — liczba sygnałów elementarnych przekazanych w czasie T_i .

Wzór ten bierze pod uwagę tzw. czynność spontaniczną.

Impulsacja spontaniczna występuje w wielu kanałach nerwowych przy braku sygnału (bodźca). Bardzo często impulsacja ta z punktu widzenia przesyłania sygnałów traktowana jest jako zakłócenie. W istocie jednak odgrywa ona ważną rolę przy odbiorze słabych sygnałów.



Ryc. 36 Reakcja pojedynczych włókien struny bębnekowej na różne roztwory NaCl (wg Gawrońskiego). Na osi odciętych — stężenie NaCl, na osi rzędnych — częstotliwość impulsacji. Włókna osmotropowe reagują na wzrost stężenia NaCl spadkiem (1) lub wzrostem (2) częstotliwości wyładowań.



Ryc. 37 Histogram interwałów wyładowań neuronów ciała kolankowatego bocznego (wg Lewisa). Na osi odciętych oznaczono wartości interwałów, na osi rzędnych liczebność występowania danych interwałów.

Jeśli przyjmiemy, że częstość impulsacji we włóknach nerwowych rośnie wprost proporcjonalnie do natężenia bodźca, według funkcji zbliżonej do logarytmicznej, wtedy łatwo sobie wyobrazić, że dla bardzo słabych bodźców częstotliwość impulsacji zbliżałaby się do zera. Mogłoby to wpływać w istotny sposób na opóźnienie odpowiedzi na słabe bodźce. Istnienie impulsacji spontanicznej i traktowanie informacji jako czynnika zmieniającego tę impulsację rozwiązuje w większości przypadków problem informacji o bardzo słabych bodźcach.

Wiele powyższych koncepcji opracowano na podstawie założeń teoretycznych wynikających z modelowania czynności neuronu w ramach prac wchodzących w zakres Moniki. Z drugiej strony istnieje szereg danych doświadczalnych potwierdzających powyższe założenia.

Na ryc. 35 do 37 przedstawiono przykłady kodowania informacji w neuronach.

Powyższy artykuł nie wyczermuje całości zagadnienia. Jest tylko przybliżeniem i zasygnalizowaniem złożoności tematu.